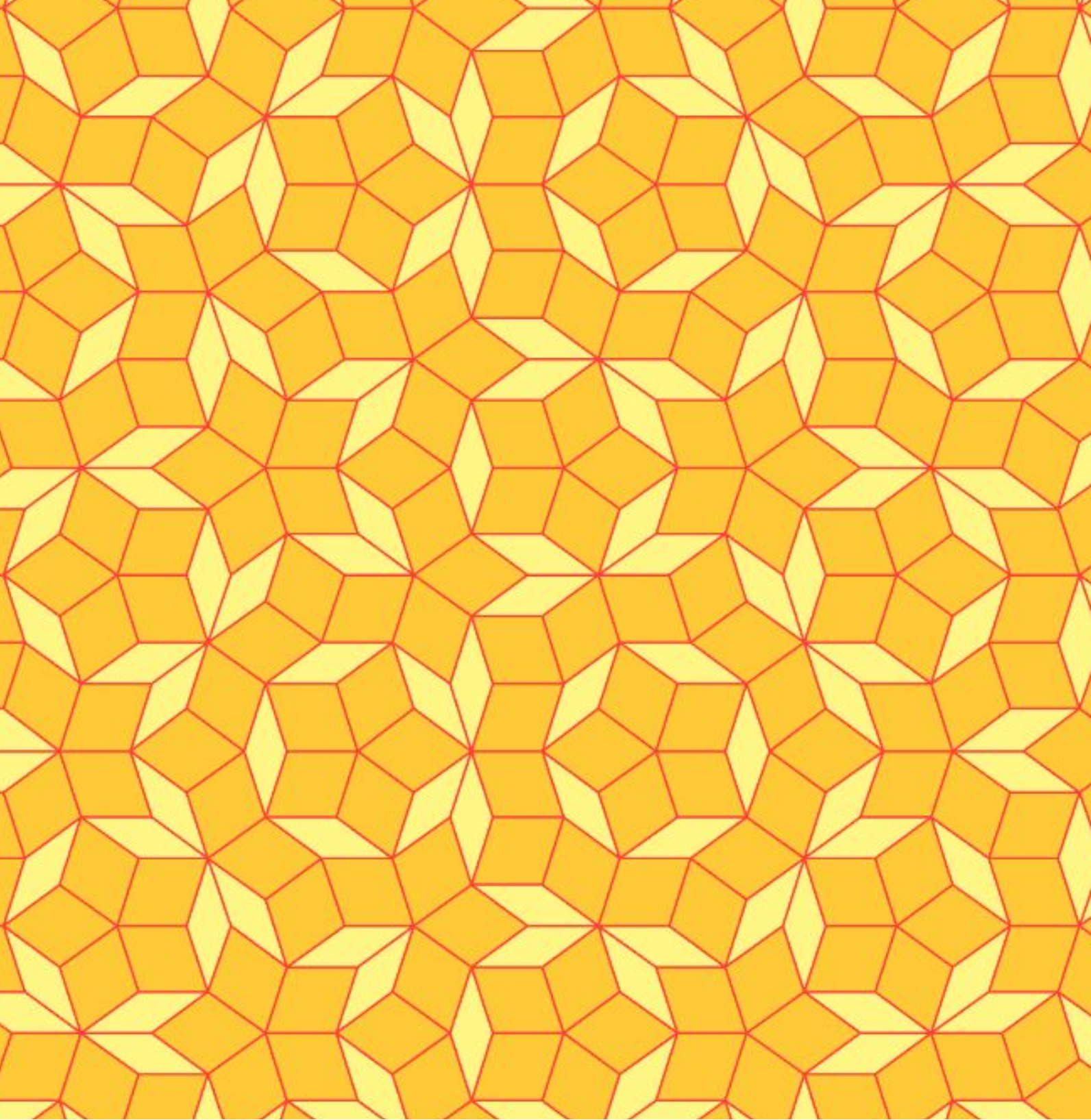




THE PENRSE MAGAZINE

STEM

Issue 9, 29 June 2026

Welcome to the nint Issue of Penrose Magazine!

Penrose is a STEM magazine where we hope to establish a community of young people who are passionate about STEM and want to share with their peers and further their knowledge beyond the curriculum. This installment of the magazine consists of the academic articles from the paradox of Schrödinger's car to the photoelectric effect. We hope to continue fostering an environment where people are encouraged to push themselves to create meaningful work and support each other to grow.

Table of Contents

3-4	The Schizophrenia Blind Spot by Ada Belkin
5-6	The Paradox of Schrödinger's Cat by Asia Pansera
7-9	From Aldehydes to Medicine: The Unnoticed Role of Chemical Synthesis in Drug Discovery by Chiamanda Christian-Iwuagwu
10-12	How Far We Are From Successfully 3d-Bioprinting Complex Human Organs For Transplants by Evina Jinsy-Benny
13-14	Abstraction, Decomposition and Their Limits by Francisca Awe
15-16	How Brain Waves Improve Cognitive Health While Asleep by Isla Heckman
17-20	Modeling Behavioral Interaction Features in AI-Driven Learning Systems by Iulia Georgiana Buhai
21-24	Chemiosmosis and Energy Production: Rotation Is Not the Only Way by Jacopo Medici
25-26	Underwater Drone Engineering by Katherine Hatcher
27-29	The Rise Of AI In The Operating Theatre By Kiana Oliver
30-33	Predicting the Invisible: AI in Non-Coding DNA Research By Lakshmi N Gowda
34-36	DNA Phenotyping: Applications, Limitations, and Ethical Challenges By Madison Krause

- 37-38 The Ghost in the Machine: Learning to Live with a Prosthetic Limb**
by Maimouna Tougoutcho Coulibaly
- 39-41 Proton, but Heavier: New Subatomic Particle Discovered in The Large Hadron Collider**
by Maja Madej
- 42-44 Investigating the Fractured Self: The Neuropsychology of Dissociative Identity Disorder.**
by Maria Wahba
- 45-47 Clot to Trot: Longevity of Localised NOSA in Artificial Placentas Across a Variety Of Genotypes in Lamb Models.**
by Marshall Zhang
- 48-49 Dirac Sea and Quantum Brain Dynamics: A Science Based Insight of Psychic Abilities**
by Omar Faruque
- 50-51 The Role of Red Light Therapy in the Aging Process**
by Pepper Koh
- 52-54 The Return of The Dire Wolf**
by Rim Laasri
- 55-57 Bacterial Vaginosis Treatments and How They Can Change Vaginal Microbiota Composition**
by Saadia Jimenez-Ñeco
- 58-60 The Role of Epinephrine in Anaphylaxis**
by Sara Glaesser
- 61-64 Photoelectric Effect And Its Practical Role in AI And Hardware**
By Weronika Padjasek



The Schizophrenia Blind Spot

No individuals with congenital blindness have ever been diagnosed with schizophrenia. This severe mental health condition, characterised by symptoms such as persistent hallucinations, delusions, disorganised thinking and anhedonia, affects approximately 0.43% of adults worldwide [1]; yet, despite over 23 million cases, the disorder is often misunderstood and thus commonly stigmatised [1].

Although the causes of schizophrenia are yet to be completely determined, scientists believe the correlation between the absence of schizophrenia and congenital blindness to be especially unusual, as early blindness often results from infections, genetic mutations or brain trauma - factors which all individually indicate a higher risk of developing psychotic disorders [2]. This anomaly may provide insights into understanding schizophrenia, which could help inform mental health services' management of this complex condition and enable more effective intervention methods to be found.

The strongest evidence demonstrating this correlation emerges from a study conducted in Western Australia in 2018, with a large sample size of 467,945 children who were born between 1980-2001 [3]. It found that of the 1,870 children who later developed schizophrenia, none were born blind [3]. One possible explanation proposed by researchers is that those with congenital blindness may possess a protective mechanism which prevents schizophrenia

from developing, although early blindness does not seem to affect the likelihood of developing neurodevelopmental disorders in general. For example, studies have suggested that blind children may show characteristics of infantile autism [2]. Consequently, several hypotheses have aimed to explain what characteristics of schizophrenia may allow congenital blindness to exhibit a 'protective' effect.

One possible hypothesis, advocated by Dr Thomas Sedlak, assistant professor of psychiatry and behavioural sciences at Johns Hopkins Medicine, is that congenitally blind people may have certain enhanced cognitive functions due to their 'brain learn[ing] to use [the] extra cortex in ways that are protective' [4]. As a result of the individual being visually impaired from an early age, their brain may compensate by developing heightened sensitivities in other areas, involving enhancement of auditory, olfactory and tactile senses [5].

These adaptations of the brain are especially prevalent in those born with blindness, as neuroplasticity is greatest at a very young age, when there is the greatest potential for new neural connections to be formed [5]. Thus, this formation of alternative neural processing pathways may be responsible for mitigating the development of certain symptoms of schizophrenia, particularly those relating to cognitive dysfunctions, such as impaired memory skills and difficulty processing the sounds and sources of speech [2]. For instance, patients with schizophrenia often

struggle to determine the source of the voice they hear, contributing to delusions [6]. In contrast, congenitally blind people are more receptive to sound localisation [5], possibly preventing them from developing certain symptoms associated with schizophrenia.

An alternative hypothesis suggests that congenitally blind people may be protected by an inability to receive false visual cues associated with the psychotic symptoms of schizophrenia [7]. Before developing psychotic symptoms, patients have been observed to have visual abnormalities such as retinal issues, unusual eye movements, and abnormal blinking rates [7]. These abnormalities lead to confusing signals about the external environment being received by the brain, forcing the individual to fill in gaps and make often inaccurate predictions about the world around them [7].

Congenitally blind people may be unable to detect inconsistencies between their knowledge from lived experience and the sensory information being received in that current moment, causing the individual to fabricate false inferences; many scientists believe this disturbance in perceptions could cause hallucinations [7]. Conversely, scientists hypothesise that when someone

is blind from birth, they rely heavily on non-visual cues to navigate their environment, thus conditioning the brain to have certain behavioural and neurophysiological adaptations [7]. For example, it has been found that congenitally blind people have an enhanced NMDA receptor function, associated with the possible higher presence of NR2B subunits, which are integral to synaptic plasticity and learning, and are primarily found in the forebrain. This results in heightened modulatory signals. These can enhance internal prediction precision, leading to a lower likelihood of the development of psychotic symptoms associated with schizophrenia [7].

Ultimately, although definitive answers regarding this phenomenon are yet to be uncovered, several practical implications emerge from this research. Firstly, findings could signify that visual and cognitive training at a young age, centred around augmenting sensory perception, attention, and memory, could be a critical method of early intervention for people at high risk for the condition [4]. Furthermore, the research could help individuals already struggling with the disorder manage symptoms through enhancing and focusing on senses other than vision [5], [7].

By Ada Belkin

References:

- [1] World Health Organisation, "[Schizophrenia](#)," Jan. 10, 2022.
- [2] E. Leivada and C. Boeckx, "Schizophrenia and Cortical Blindness: Protective Effects and Implications for Language," *Frontiers in Human Neuroscience*, vol. 8, Nov. 2014.
- [3] V. A. Morgan, et al., "Congenital Blindness is Protective for Schizophrenia and Other Psychotic Illness. A Whole-Population Study," *Schizophrenia Research*, vol. 202, pp. 414–416, Dec. 2018.
- [4] J. Tzeses, "[The Curious Link between Blindness and Schizophrenia](#)," *HealthCentral*, Mar. 21, 2020.
- [5] D.-R. Chebat, F. C. Schneider, and M. Ptito, "Spatial Competence and Brain Plasticity in Congenital Blindness via Sensory Substitution Devices," *Frontiers in Neuroscience*, vol. 14, 2020.
- [6] NHS, "[Symptoms – Schizophrenia](#)," 2023.
- [7] T. A. Pollak and P. R. Corlett, "Blindness, Psychosis, and the Visual Construction of the World," *Schizophrenia Bulletin*, vol. 46, no. 6, 2019.



The Paradox of Schrödinger's Cat

The concept of a cat being both alive and dead at the same time may seem unsettling, but according to quantum mechanics such a possibility would not be considered contradictory. Instead, both outcomes are both paradoxical and plausible simultaneously.

At the beginning of the 20th century, physicists believed that classical physics could only explain the behaviour of macroscopic matter. However, when scientists attempted to explain atomic-scale phenomena and particles, such as electrons, protons and neutrons, using the same laws, those laws no longer seemed to apply. To amend this issue, many physicists of the 1920s, including Max Planck, Albert Einstein, Niels Bohr and Werner Heisenberg developed new theories. These theories transformed the understanding of physics from a deterministic one to a probabilistic one, giving rise to the field of quantum mechanics.

Central to quantum mechanics is the idea that particles exhibit wave-like and particle-like behaviour depending on how they are measured. Wave-like behaviour allows molecules to exist in a superposition of multiple states, whereas particle-like behaviour describes molecules as discrete units. As a result, systems are described in terms of probabilities instead of certainties. One of the most important theoretical physicists and central figures in the development of quantum mechanics is Erwin Schrödinger, known for having

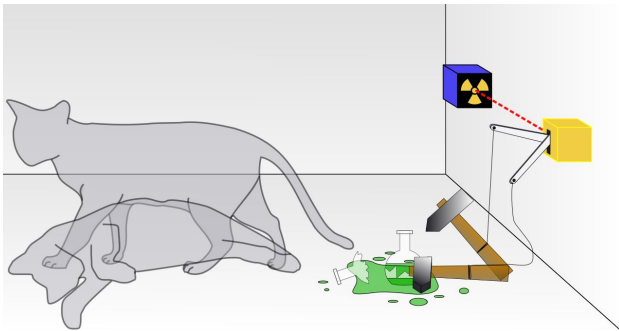
formulated the “*Schrödinger equation*”. This equation describes how the quantum state of a physical system evolves over time. It is the equivalent of Newton’s second law:

$$F = ma$$

Just as Newton’s second law governs how a classical object’s motion evolves under a force, the “*Schrödinger equation*” governs how the wave function, the fundamental variable in the mathematical expression, of a quantum system evolves under the influence of energy. This mathematical tool is usually denoted as “ Ψ ”, and it represents all the possible states of a system. “ Ψ ” itself carries no direct physical meaning; it is its absolute square, $|\Psi|^2$, that yields observable information. This quantity defines the probability density of locating an electron within a specific region of space at a given time [1]. In 1935, Erwin Schrödinger proposed a thought experiment meant to illustrate the absurdity of the Copenhagen interpretation of quantum mechanics. According to this interpretation, a system, even a macroscopic one, exists in a “superposition” of possible states until it is observed or measured, forcing the quantum states into one of the possible outcomes [2].

Not every physicist agreed with the Copenhagen interpretation. Some considered the Many-Worlds Interpretation (MWI) more logical, as it claims that all possible outcomes of an event actually occur, but each in its own separate branch of the universe [1], [3]. Nevertheless, the MWI remains controversial; since these

parallel universes cannot be observed or tested, many physicists argue it cannot be considered a scientifically provable theory. To challenge such interpretations, Erwin Schrödinger designed a thought experiment, known as "The paradox of Schrödinger's cat", where a cat gets sealed in a box with a small quantity of a radioactive substance whose atoms have a 50% chance of decaying within a given time. If the supposed event occurs, a hammer gets activated to break a vial with cyanide inside, leading to the cat's death [4]. According to the Copenhagen interpretation, before the observer interacts with the box and the cat, it exists in a quantum superposition of both alive and dead states simultaneously. These two outcomes coexist until an observer interacts with the system, at which point the act of measurement forces it to resolve into one definitive state.



Schrödinger's point was not about a lack of knowledge of the cat's state, but rather to illustrate and criticize the claim of the Copenhagen interpretation. The paradox is meant to underline two controversial aspects. The first is physical: an indeterminacy in the microscopic world generates paradoxical consequences in the macroscopic one, since quantum superposition, where a particle can exist in

multiple states simultaneously, does not translate intuitively to the large-scale objects we observe in everyday life. The second is philosophical, raising the deep question of whether reality exists independently of observation, or whether measurement plays a fundamental role in defining it.

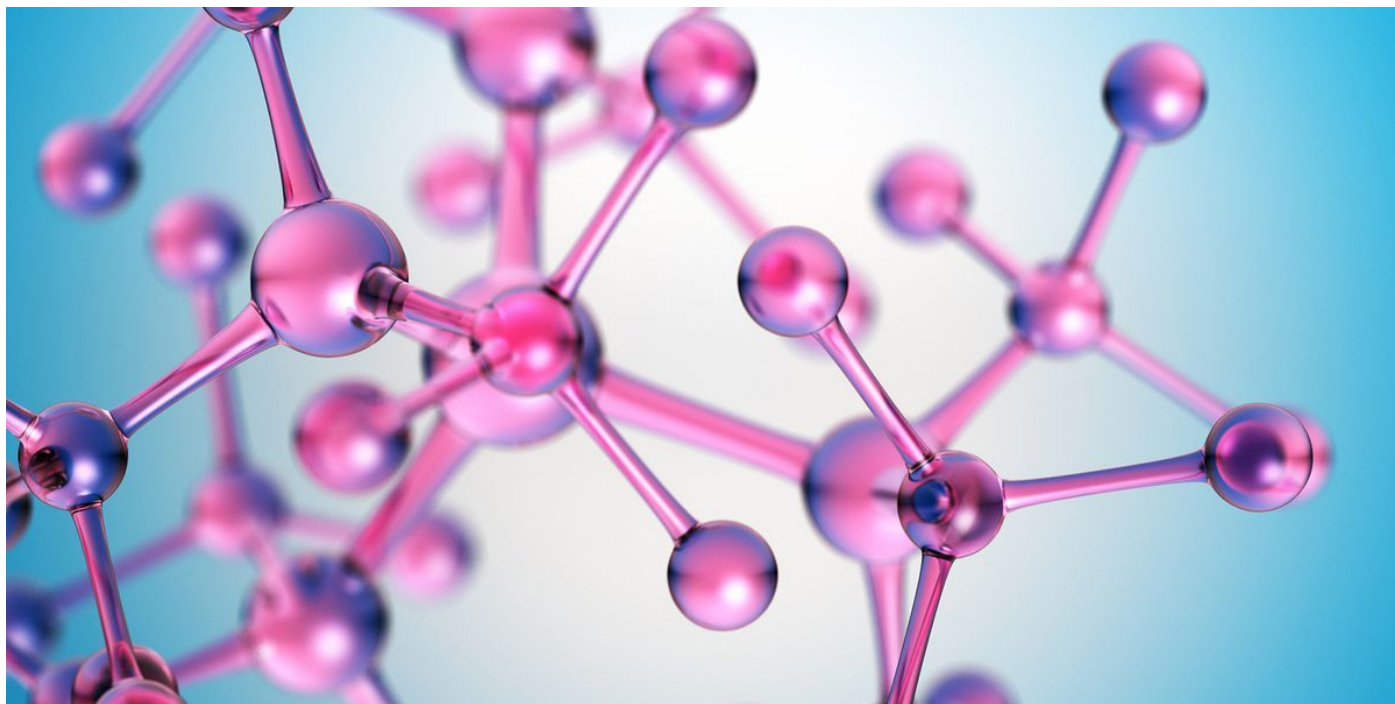
The cat, simply by being inside the box, is constantly in contact with air molecules, radiation and the container itself that effectively "decohere" the system and cause, as we mentioned, a theoretical macroscopic superposition to become unobservable. This phenomenon has been called "quantum decoherence" and it represents the modern perspective proposed by physicists when trying to provide an explanation for the Schrödinger paradox. The core idea is that the bigger a system is and the more it interacts with its surrounding environment, the less observable and stable its superposition of states becomes [5], [6].

Erwin Schrödinger did not come up with his paradox to be taken literally. However, nearly a century after it was proposed, the paradox still portrays the powerful and strange nature of how quantum mechanics challenges basic intuitive understanding about reality. At the turn of the twentieth century, they were firmly convinced they had all the laws of physics figured out and were close to a complete description of nature, only to be defied by a cat that refused to be merely alive or dead. This seemingly playful thought experiment forced a confrontation with the very limits of how reality itself can be defined [7].

By Asia Pansera

References:

- [1] Stanford Encyclopedia of Philosophy, "[Many-Worlds Interpretation of Quantum Mechanics](#)."
- [2] H. D. Zeh, "[On the interpretation of measurement in quantum theory](#)," arXiv, Apr. 4, 1999.
- [3] H. Everett III, "[Relative State Formulation of Quantum Mechanics](#)," arXiv, Feb. 2, 2001.
- [4] Quantum Girl, "[IL PARADOSSO che ha ROTTO la FISICA: il GATTO di SCHRÖDINGER](#)."
- [5] Stanford Encyclopedia of Philosophy, "[Decoherence and the quantum-to-classical transition](#)."
- [6] W. H. Zurek, "[Decoherence and the transition from quantum to classical](#)," arXiv, Jun. 10, 2003.
- [7] V. Benzi, "[La scienza delle meraviglie](#)."



From Aldehydes to Medicine: The Unnoticed Role of Chemical Synthesis in Drug Discovery

The Fundamental Chemistry of Aldehydes

Aldehydes are one of the most important and widely encountered organic compounds, playing a significant role in nature and industry. Characterized by its carbonyl group bonded to hydrogen, they hold distinctive properties such as strong odors, high boiling points and a relatively high chemical reactivity [1]. They can be produced from alcohols and are widely utilized in industrial, biological, and pharmaceutical applications [1].

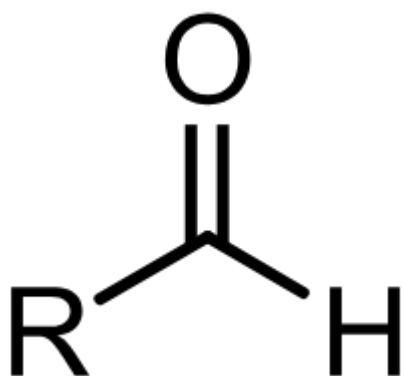


Fig. 1. The representative structural formula of aldehydes (R-CHO).

A salient feature of aldehydes that makes them invaluable in developing drugs is their pronounced reactivity and versatility as synthetic intermediates [2]. Due to the

strong polarization of the carbonyl group, the carbonyl carbon is highly electrophilic and readily undergoes nucleophilic addition reactions. This enables the formation of new carbon-carbon bonds and facilitates the synthesis of complex molecular structures [2]. It is noteworthy that over 30% of pharmaceutical intermediates are derived from aldehyde derivatives [3].

Innovations in Aldehyde Functionalisation

One crucial contribution to drug synthesis using aldehydes is the light-powered approach developed by the researchers at the University of Hawai'i at Mānoa, that directs aldehydes through highly controlled chemical reactions. It is a sustainable approach, producing significant molecular structures used in drug development, chemical manufacturing, and product study [4]. Furthermore, aldehyde-based synthesis methods contribute to sustainable chemistry by creating efficient reaction pathways that require fewer synthetic steps, generate less chemical waste, and improve atom economy. These advantages help to reduce the environmental impact of chemical production while still maintaining the versatility required for the synthesis of complex molecules.

Deoxygenative Difunctionalization

Recent advances in synthetic chemistry are

transforming how scientists approach drug discovery, underlining a crucial role for chemistry in innovation and healthcare that is frequently overlooked [5]. The deoxygenative difunctionalization of aldehydes enables the transformation of simple and abundant carbonyl feedstocks into structurally complex, high-value molecules [5]. Researchers at the University of Hawai'i at Mānoa developed a visible-light-driven, redox-neutral system that generates ketyl radicals from α -acyloxy halides and merges this intricate process with hybrid palladium (Pd) catalysis to enable controlled and efficient bond construction [5]. In this reaction, ketyl radicals act as highly reactive intermediates that facilitate the formation of new carbon-carbon bonds, enabling the efficient synthesis of complex molecular structures relevant to drug development and discovery.

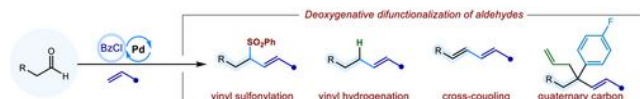


Fig. 2. An illustration of the deoxygenative difunctionalization, where an aldehyde reacts under palladium catalysis to form multiple products [4].

This strategy allows aldehydes to function as adaptable one-carbon building units, which produces allylic sulfones when nucleophiles are present and conjugated dienes when nucleophiles are absent [2], [5]. The method is particularly versatile, displaying effectiveness across a vast range of vinylarenes, aldehydes, and sulfonyl compounds, and even has the ability to modify complex natural products and pharmaceuticals at later stages of synthesis [5]. The empirical value of the products is demonstrated through gram-scale reactions, regioselective $SN2'$ substitution reactions enabling the formation of quaternary carbon centers, and a one-pot reductive hydrovinylation sequence [5]. By integrating ketyl radical chemistry with hybrid palladium catalysis, this work demonstrates that aldehyde deoxygenative difunctionalization is an effective and reliable method for constructing complex molecular structures [5].

Potential Impact in Pharmaceutical Development

Although the light driven deoxygenative difunctionalization method has not yet been used in drug manufacturing, its potential impact on pharmaceuticals is substantial [5], particularly for structurally complex drug candidates. Drugs such as Paclitaxel and Docetaxel, which are chemotherapy medications used to treat a wide range of cancers, require highly intricate carbon molecules that are challenging to construct using traditional synthetic routes [6]. Likewise, antiviral agents such as Remdesivir, which inhibits viral replication and is used in the treatment of COVID-19, rely on detailed molecular modification to maintain efficiency. Similarly, central nervous system drugs such as Paroxetine, a selective serotonin reuptake inhibitor (SSRI) commonly prescribed as an antidepressant, often contain structural motifs, including quaternary carbon centers that present considerable synthetic challenges [6].

The capacity of this method to efficiently form carbon-carbon bonds, enable late-stage functionalization, and construct complex molecular structures underscores its relevance to modern drug development. Late-stage functionalization refers to the selective chemical modification of advanced intermediate drug candidates, allowing their properties to be turned without resigning the entire synthetic route. This is particularly valuable in pharmaceutical synthesis, as it enables rapid optimization of biological activity, solubility, and safety while minimizing the need for extensive resynthesis of complex molecules [7]. This reinforces the argument that, despite the emphasis on biological targets in drug discovery, it is the underlying chemistry that precisely and accurately determines which compounds can be developed and optimised.

Perspectives from Currents Research

The University of Hawai'i researchers frame this light-driven, redox-neutral aldehyde deoxygenative difunctionalization strategy

as a remarkable contribution to improving the efficiency and attainability of modern chemical synthesis, with wide-ranging implications for both societal benefit and scientific progress [4]. They highlight that the development of more sustainable and efficient methods for constructing complex molecular structures has the potential to accelerate the pharmaceutical development, reduce costs, and improve access to essential products, including therapeutics and advanced materials. Additionally, the study highlights the role of ongoing innovation in overcoming inherent limitations, thereby reframing chemistry as a catalyst for discovery rather than a constraint [4].

Limitations and Constraints

Despite its potential, the light-driven deoxygenative difunctionalization of aldehydes has several limitations, as briefly underlined initially, that restrict its immediate use in pharmaceutical synthesis. As this method is immensely dependent on specialised reagents such as α -acyloxy halides, additional preparation may be required, reducing the overall efficiency [5]. The use of palladium also raises concerns regarding cost and sustainability due to its limited availability [5], [4]. Additionally, light-driven reactions require precise and accurate conditions, making large-scale

application challenging. The involvement of reactive ketyl radicals can potentially lead to side reactions, affecting the selectivity and yield. These limitations are exceptionally relevant in the synthesis of complex drugs such as remdesivir and paclitaxel, where high precision is crucial.

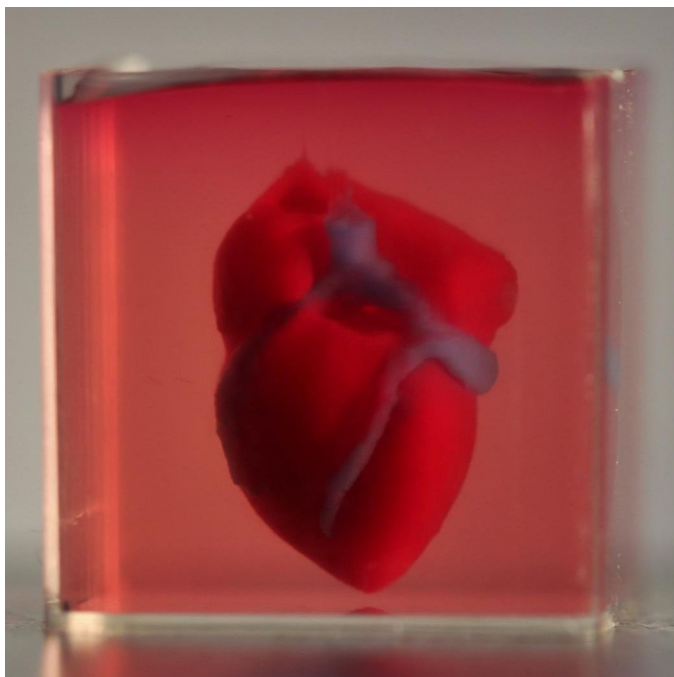
Conclusion

To conclude, the deoxygenative difunctionalization of aldehydes represents a notable development in synthetic chemistry, demonstrating the capacity to convert simple carbonyl compounds into structurally complex molecules through light-driven processes. The integration of ketyl radical chemistry with palladium catalysis enables efficient carbon-carbon bond formation and widens the range of accessible molecular architectures relevant to pharmaceutical research [5]. Its applicability to the synthesis of advanced molecules, including those present in the drugs such as Paclitaxel and Remdesivir, demonstrates its significance in modern medical chemistry [6]. Ultimately, this work supports the conclusion that advances in chemical synthesis play a determinative role in shaping and enabling drug discovery, highlighting the underlying importance in the optimisation of pharmaceuticals and drug discovery.

By Chiamanda Christian-Iwuagwu

References:

- [1] EBSCO Information Services, Inc., "Aldehydes," 2022.
- [2] Khan Academy India – English, "Mechanism of nucleophilic addition reactions | Aldehydes, ketones & acids | Chemistry | Khan Academy," YouTube, Jun. 29, 2024.
- [3] ReAnIn, "Aldehydes Market Size & Industry Share," Reanin.com, Mar. 2026.
- [4] University of Hawai'i System News, "New chemical discovery could speed up future medicines, materials," University of Hawai'i System News, Dec. 03, 2025.
- [5] T. Bian, et al., "Deoxygenative difunctionalization of aldehydes via ketyl radical and light/dark Pd synergy," *Angewandte Chemie*, vol. 138, no. 3, Nov. 2025.
- [6] T. Cernak, et al., "The medicinal chemist's toolbox for late-stage functionalization of drug-like molecules," *Chemical Society Reviews*, vol. 45, no. 3, pp. 546–576, 2016.
- [7] A. P. Montgomery, et al., "An update on late-stage functionalization in today's drug discovery," *Expert Opinion on Drug Discovery*, vol. 18, no. 6, pp. 597–613, Apr. 2023.



How Far We Are From Successfully 3d-Bioprinting Complex Human Organs For Transplants

Recently, researchers at Tel Aviv University achieved a world first: a miniature, 3D printed heart complete with its own blood vessels. While only the size of a rabbit's heart, this experiment advances in the medical world, continuing the new era in regenerative medicine [1].

3D bioprinting is an additive manufacturing technology that uses biomaterials or biocompatible materials to create 3D tissues and other constructs [2]. This process dates back to 1988, where Robert Klebe used a method called cytoscribing. He modified an inkjet printer to deposit fibronectin patterns onto collagen substrates and seeded cells onto these protein templates. By doing this, and stacking them together strategically, Klebe created the first ever form of simple 3D tissues [3].

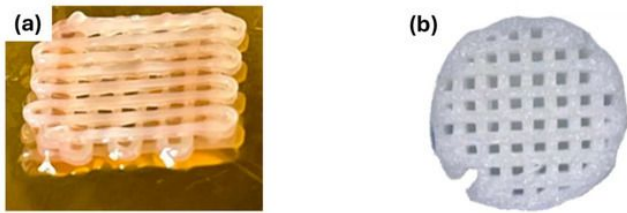
Later in 1996, Ramsey Foty and his team, found that if cells are placed in close proximity, they can fuse and become tissues without needing a solid scaffold to support the cells structurally. This idea showed that embryonic cells can behave like liquid droplets and will reassemble due to surface tension and differential cell adhesion [4].

In 2003, Chris Wilson and his colleagues modified a standard Hewlett-Packard inkjet printer to dispense living cells as “ink,” creating the first proper 3D bioprinter. Inkjet-based bioprinting operates on the same principles as traditional inkjet printing, but instead of ink, it uses bioink-containing living cells [5], [6]. Using this method, they printed both biomolecules and viable cells in predefined patterns. These printed endothelial cells remained alive, and the patterns remained intact; however, one problem that is faced during inkjet-based printing is the bioink viscosity requirement, which in simplified terms means printing certain cell densities would be difficult and take a very long time to build larger cell structures.



Subsequently, in 2004 a scientist named Gabor Forgacs used this same “ink” method and created tubes of sheets of cells, and through this he formed a way for cells to biologically fuse together without the need of using scaffolding internally creating tissues. This work laid the groundwork of his later discoveries at his very own established company under the name “Organovo”; they crafted the first 3D bioprinted human liver tissue which was later sent for drug testing. This was brought upon by Computer-Aided Design (CAD)

software, which was used to design the models of the tissues digitally, and to guide the bioprinter in constructing the desired tissues needed [6].

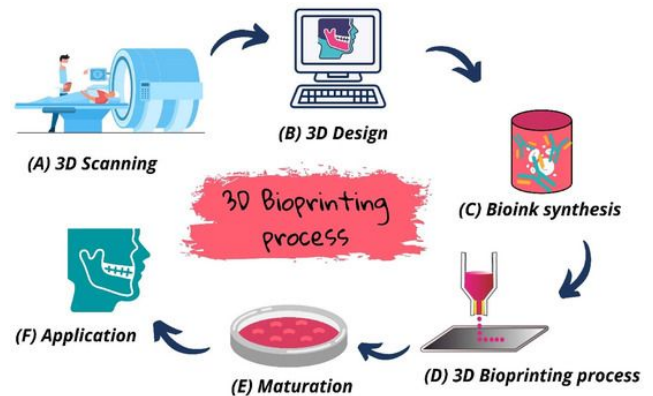


The first 3D bioprinted human liver tissue was created using a research series of clinical trials conducted at Boston Children's Hospital. Here, hand built replacement bladders were constructed from scaffolds of collagen and synthetic polymer. These were then stacked with cells from seven patients, allowing them to grow into functioning organs. Seven years after the replacement bladders were implanted, all patients were still in good health, but the process of constructing the scaffolds was very strenuous [7].

As a result, in 2006, Anthony Atala, one of the members of researchers that had conducted those trials, and his team at the Wake Forest Institute for Regenerative Medicine (WFIRM) found a solution and successfully printed a miniature human kidney using 3D bioprinting. Atala and his team demonstrated an early-stage experiment that could someday solve the organ donor problem: a 3D printer that uses living cells (bio-ink) to output a transplantable kidney. Although the bio-ink replicated the kidney tissue, the tissues were not living [8].

Researchers at Beijing's Tsinghua University developed a micro-bioprinting platform designed to enter the body via an endoscope to repair gastric wounds in real-time. This was achieved by testing the microbot and the delivery system with a biological model of a human stomach and an endoscope to mimic the insertion and bioprinting operation. Additionally, they also carried out a bioprinting test in a cell culture dish to test how effective the device was at bioprinting viable cells and repairing wounds. Despite the high viability and steady proliferation of this test, Professor

Tao Xu stated that further work is necessary, including reducing the size of the bioprinting platform for easier maneuvering and developing more effective bioinks [9].



Despite these setbacks, in 2019, researchers at Tel Aviv University created the first 3D vascularized, engineered heart using a patient's own cells and biological materials. This method enables the fabrication of artificial tissue to be embedded with its own complex, branching networks of blood vessels to ensure that it receives the vital nutrients and oxygen it needs to function. For the research, a biopsy of fatty tissue was taken from patients. These cells were reprogrammed to become pluripotent stem cells, a type of stem cell that has the ability to divide and develop into almost any cell type in the body. During this process, extracellular matrix (ECM), a three-dimensional network of extracellular macromolecules such as collagen and glycoproteins, were processed into a personalized hydrogel that served as the printing "ink." After being mixed with the hydrogel, the cells were efficiently differentiated to cardiac or endothelial cells. This led them to create patient-specific, immune-compatible cardiac patches with blood vessels and, subsequently, an entire heart.

Previous attempts at 3D-printing a heart have succeeded in creating its structure, but these did not include cells or blood vessels. The team's findings indicate their method could be used to engineer custom tissue and organ replacements in the future [10]. Although the developed heart is very small and has not yet been tested in clinical trials, the team plans to develop the ability for the printed cells to form a functioning

pump and to test these hearts in animal models to eventually replace damaged human hearts.

In 2022, a shift from theoretical bio-engineering to practical, clinical application had occurred when a technique to successfully create 3D printed living tissue ear transplant was created. This was achieved using a patient's own chondrocytes (cartilage) cells. Using 3DBio Therapeutics, they provided a solution to a challenge that plastic surgeons had faced for years regarding external ear reconstruction [11].

The impact of 3D-printed organs being produced would be pivotal. In the UK alone,

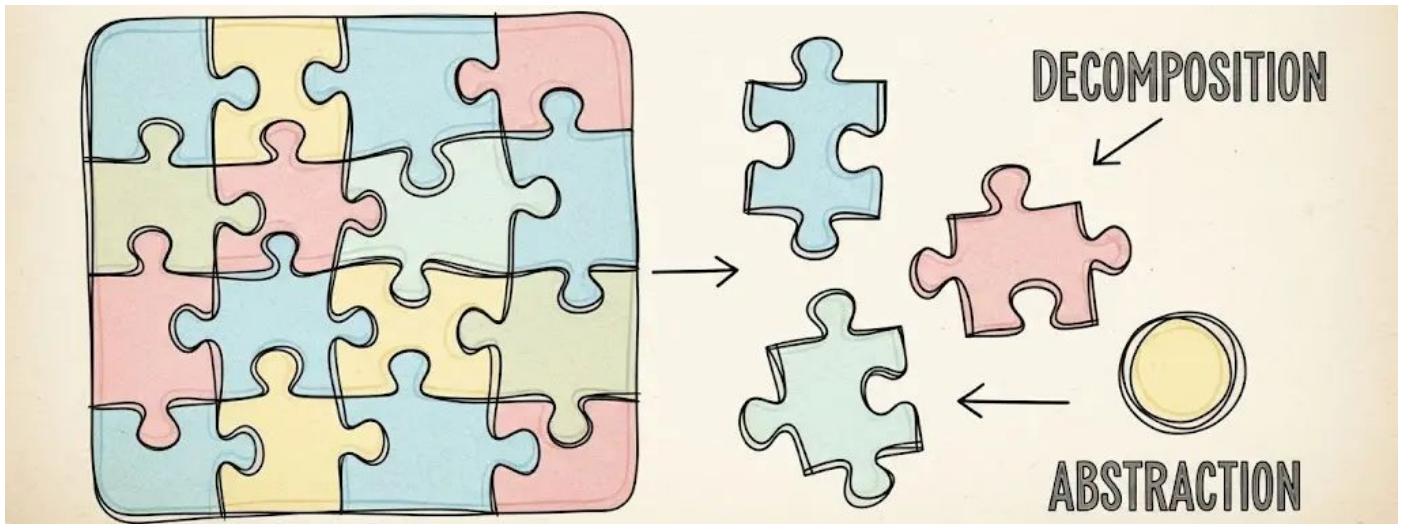
approximately 8,000 people are currently on the active transplant waiting list, where less than 60% receive a transplant, with over 460 people passing away while waiting for a transplant just last year. Even when organs become available, the risk of rejection remains high, leading to other risks and dangers that a patient must face [12].

Bioprinting could help in two ways: expanding the supply of transplantable organs and allowing organs to be made using a patient's own cells, reducing the risk of rejection substantially. This combination of availability and compatibility would mark a revolution in both regenerative and transplant medicine.

By Evina Jinsy-Benny

References:

- [1] N. Noor, A. Shapira, R. Edri, et al., "3D Printing of Personalized Thick and Perfusable Cardiac Patches and Hearts," *Adv. Sci.* 2019.
- [2] "What is 3d bioprinting?," Cellink by Bico.
- [3] R. J. Klebe, "Cytoscribing: A Method for Micropositioning Cells and the Construction of Two- and Three-Dimensional Synthetic Tissues," *Exp. Cell Res.* 1988, 179, 362–373.
- [4] R. A. Foty, C. M. Pflieger, G. Forgacs, M. S. Steinberg, "Surface Tensions of Embryonic Tissues Predict Their Mutual Envelopment Behavior," *Development*, 1996.
- [5] "Inkjet bioprinting of biomaterials," *Chem. Rev.* 19, 2020.
- [6] A. Skardal, D. Mack, E. Kapetanovic, et al., "Bioprinted Amniotic Fluid-Derived Stem Cells Accelerate Healing of Large Skin Wounds," *Stem Cells Transl. Med.*, 2012.
- [7] M. Whitaker, "The history of 3d printing in healthcare," *Jul.* 2014.
- [8] "3D Printing: The Answer to a Shortage of Human Organs?," *Science World*, Oct. 22, 2013.
- [9] G. Ying, Xinhuanet.com "New bioprinting method for gastric wounds."
- [10] "TAU scientists print first ever 3D heart using patient's own cells," *Tel Aviv University*.
- [11] Y. Ma, M. S. Lloyd, "Systematic review of Medpor versus autologous ear reconstruction," *J Craniofac Surg.* 2022;33:602–606.
- [12] "Organ transplant waiting list hits record high as donor and transplant numbers fall," *NHS Blood and Transplant*, Jul. 9, 2025.



Abstraction, Decomposition and Their Limits

Introduction

Abstraction and decomposition are both effective methods used to solve problems in computer science and other fields. Both of these methods simplify complex problems and allow for focus on important aspects of problems; they have been proven useful for many types of problems, but whether they are always helpful is a topic of debate.

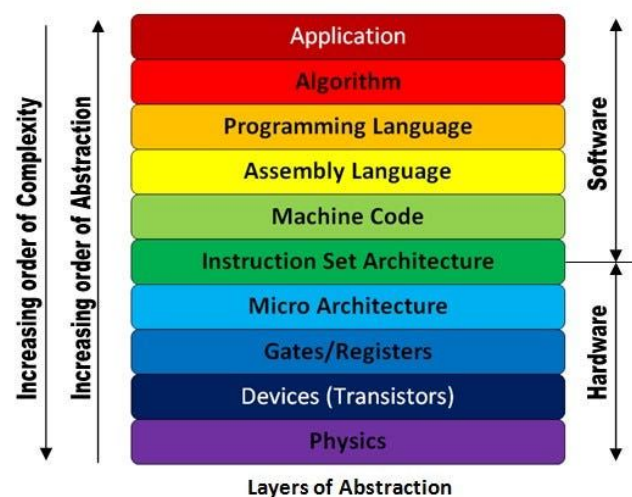
An Overview of Abstraction and Decomposition

Abstraction is derived from the Latin *abs*, meaning *away from* and *trahere*, meaning *to draw*. It is the process of taking away or removing characteristics from something in order to reduce it to a set of essential characteristics [1]. Decomposition in computational thinking is the process of breaking down a problem into a number of smaller problems that can more easily be addressed [2]. Abstraction and decomposition are often used together to solve complex problems. When breaking a problem down into sub-problems, each sub-problem is abstracted from the others, this results in the sub-problems being mostly independent of each other.

Abstraction and Decomposition in Action

Neural networks use decomposition in pattern recognition. For example, to recognise an image of a cat, the network will need to identify the main characteristics of a cat. To identify these features, the

network will break each feature down into smaller parts, with no limit to how much it breaks them down [3].



An example of decomposition and abstraction working hand in hand is the layers of abstraction. Each layer solves a sub-problem that makes up the problem. Every layer of abstraction has been abstracted from the previous one(s), and to work with a layer, the in-depth knowledge of the previous layer(s) is not needed [4].

For example a programmer does not have to worry about understanding machine code, logic gates, instruction set architecture or how transistors work because all of these have been abstracted away [4]. Furthermore, this is the reason why a user does not need to know programming languages to use an app, all the complexities of transistors, logic gates, machine code and assembly language have been converted into intuitive interfaces for the user [4].

The Limits of Abstraction and Decomposition

Abstraction and decomposition are very effective methods but they have some limitations. The Law of Leaky Abstractions is an example of a limitation of abstraction [5]. This law states that abstractions, which are crucial in computing for simplifying interactions with complex systems, will inevitably expose some aspects of their underlying mechanisms [6]. These ‘leaks’ can lead to unexpected behaviour, necessitating a deeper understanding of the underlying systems [6]. For example, In Structured Query Language (SQL), a standardized programming language that is used to manage relational databases and perform various operations on the data in them [10], certain queries can run drastically slower than others despite being logically equivalent [6]. For example, ‘ $a=b$ and $b=c$ and $a=c$ ’ runs faster than ‘ $a=b$ and $b=c$ ’ [5]. This is due to the underlying database mechanics that have been abstracted leaking through; while abstractions simplify our interaction with complex systems, they invariably expose some of their complexities [6].

There are computationally irreducible problems where applying decomposition does not make a difference to how long it takes to solve them. Wolfram’s notion of computational irreducibility says that some systems or processes cannot be simplified or accelerated beyond their natural course.

In other words, there is no shortcut to predicting the outcome of these systems without actually simulating the entire process [7].

By breaking down a complex system into smaller components, there is a risk of over-simplifying the system and not accurately representing its complexity. Additionally, decomposition can overlook the interactions and relationships between different parts of the system as it focuses on individual components, leading to a lack of integration [8].

Conclusion

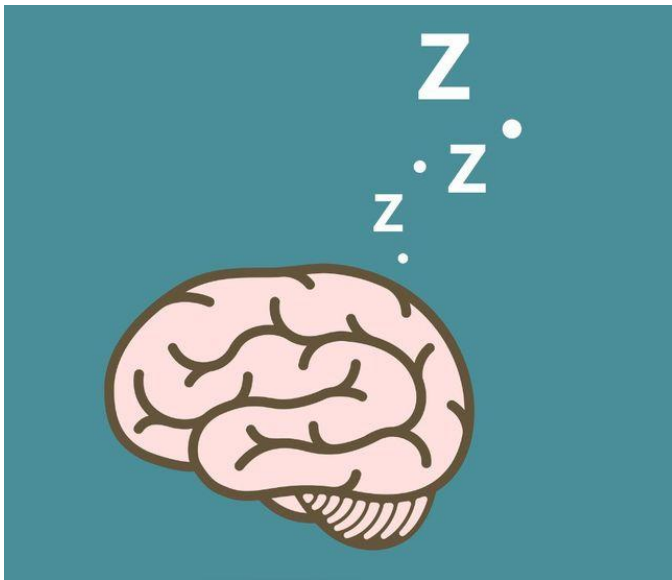
Both abstraction and decomposition are powerful and useful tools for problem solving, but just like any other tool they have their limitations. Abstractions are reductions by design, and most will eventually fail to capture something important. Even when an abstraction does not perfectly represent a system, it teaches key features and bootstraps further understanding [9].

Similarly, decompositions can fail to capture the important interactions between sub-problems leading to a disjointed solution. Although decomposition cannot speed up computationally irreducible problems it can provide structure that aids understanding. Understanding that abstraction and decomposition are not perfect, is something to keep in mind when solving problems to allow a more flexible approach.

By Francisca Awe

References:

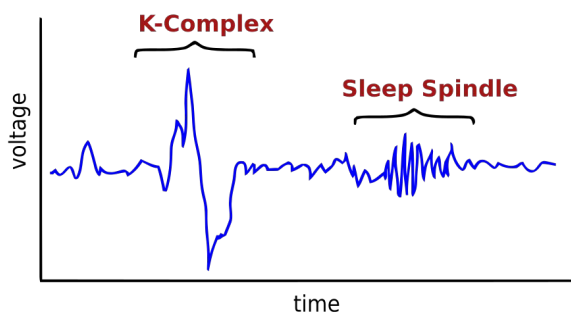
- [1] I. Wigmore, “What is abstraction?,” *WhatIs.com*, Dec. 2021.
- [2] Lcom Team, “What does decomposition mean in computational thinking,” *Learning.com*, Aug. 29, 2022.
- [3] 3Blue1Brown, “But what is a neural network? | Deep learning, chapter 1,” YouTube, Oct. 5, 2017.
- [4] BIPLOVE YT, “Every level of abstractions in computer science explained in 3 minutes,” YouTube, Oct. 25, 2025.
- [5] J. Spolsky, “The law of leaky abstractions,” *Joel on Software*, Nov. 11, 2002.
- [6] M. Egger, “The law of leaky abstractions: A guide for the pragmatic programmer,” *Medium*, Feb. 2, 2024.
- [7] Hemanth, “How to understand computational irreducibility,” *Medium*, Apr. 2, 2023.
- [8] S. Datta, “What is functional decomposition?,” *Baeldung on Computer Science*, Feb. 14, 2023.
- [9] A. Towell, “Uses and limits of abstractions,” *metafunctor.com*, Jun. 17, 2023.
- [10] P. Loshin and J. Sirkin, “What Is Structured Query Language (SQL)?,” *TechTarget*, Feb. 2022.



How Brain Waves Improve Cognitive Health While Asleep

The brain has four main stages of sleep, each of which help the brain in a different way. From consolidating memories to removing toxic proteins, the different stages of sleep each do vastly different things, but all are vital for a properly functioning brain.

During the first half of the night, humans mostly experience the three stages of non-rapid eye movement (non-REM) sleep. These stages are each characterized by different types of brain rhythms, which are synchronous neural activity regulated by the thalamus [10]. As the body experiences deeper stages of sleep, the amplitude of brain waves increases. Amplitude is the distance between the highest and lowest point on a wave. Higher amplitude brain waves slow brain function.

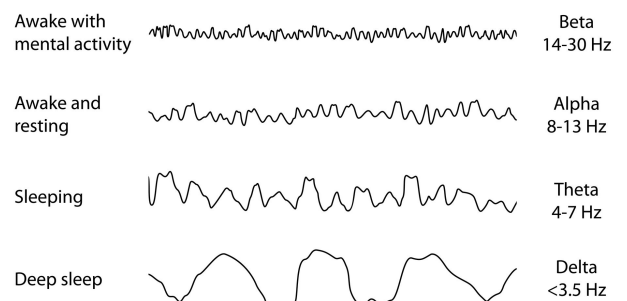


During stage one, lasting for approximately ten minutes, muscles relax, and respiration slows. This stage is characterized by the presence of theta waves: a type of brain wave that has a blend of different frequencies and low amplitude [10]. Cognitive health is largely unaffected during

this time, with the brain mainly signaling to the body to release stress. It is very easy to be woken up during stage one, and the brain can still process external stimuli if required [6].

The second stage primarily functions to keep the body asleep [5]. It does this through two key types of brain activity. The first is K complexes: a sudden large increase in electrical activity in neurons, followed by a sudden decrease [5]. The increase happens when stimuli reaches the thalamus, and the decrease happens as the message is sent to the cerebral cortex. This ensures the brain will not be woken up by trying to process information.

Sleep spindles are continuous short bursts, 0.5 - 3 second, of high neural activity [5]. Similar to K complexes, sleep spindles block information from reaching the cerebral cortex. Typically, sleep spindles follow closely after K complexes. Additionally, the repetitive activation of neurons through K complexes and sleep spindles helps strengthen and prepare the brain for deep sleep [10].



Around an hour after falling asleep, the brain enters stage three, also known as deep sleep, and experiences delta waves [10]. This is when the brain improves and repairs itself. Neurons are improved through myelination: a process that allows information to travel through the neuron faster. It does this by encasing part of the neuron in fatty tissue. The cells known as oligodendrocytes produce the fatty tissue used for myelination [4]. These myelinated neurons, made of white matter, play a significant role in cognitive health.

The brain is also repaired during this stage through the removal of toxic materials.

Cerebrospinal fluid is constantly pulsing in and out of the brain. During deep sleep these pulses synchronize with brain waves and flush out a toxic protein called beta-amyloid [1]. This protein builds up in the brain during the day through the use of transmembrane proteins in neurons [1]. When the brain is not able to remove beta-amyloid during sleep, it has to compensate while awake. If the cerebrospinal fluid has to continue pulsing during wakefulness, it can cause reduced cognition, commonly referred to as brain fog [1]. A buildup of this protein is correlated with Alzheimer's disease and early cognitive decline.

Brain waves during REM sleep are much faster than non-REM sleep. This is when most dreams occur. Brain waves during REM sleep are called ponto-geniculo-occipital (PGO) waves. PGO waves originate in the brain stem and continuously stimulate neurons [9]. This stimulation causes dreams, which improve memory and emotional wellbeing [8].

As the brain is able to revisit memories while feeling less emotionally strained, emotions are more effectively processed during dreams. A decrease in the activity of the amygdala plays a major role in emotional recuperation, as it is a part of the brain that often stimulates fear and anxiety [2].

Lastly, while dreaming in REM sleep, the body enters a type of paralysis called

atonia. Atonia is caused by hyperpolarization of motor neurons [7]. Without this, everyone would physically act out their dreams, which often leads to injury. Fortunately, hyperpolarization of neurons makes it extremely difficult for signals to be sent to the muscles. Additionally, neurotransmitters that help with movement, including norepinephrine, serotonin, and histamine, are barely present in motor neurons during this time [7].

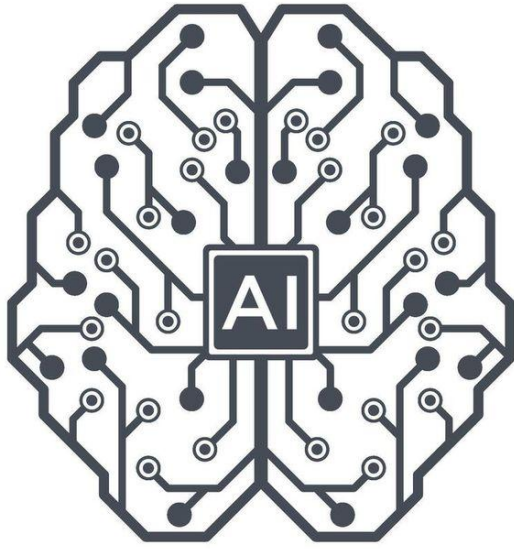
The medulla, a part of the brain that controls heartbeat, breathing, and blood pressure, also blocks muscle movement during this time. It does this by increasing inhibitory neurotransmitters like GABA and glycine near the spinal cord [6]. The medulla is able to do this because it's the part of the brain closest to the spinal cord [3]. This combination results in very little movement during REM sleep aside from faint twitches, vital muscle movements, and eye movements.

Without sleep, and the critical contributions of each sleep stage, cognitive function deteriorates rapidly. Across the four stages of sleep the brain regulates sensory input, consolidates memories, removes toxic proteins, and restores emotional wellbeing. While some of this work continues during wakefulness, it is far less effective without the synchronized activity and reduced external demands that occur during sleep. These processes therefore make sleep essential for maintaining focus, long-term cognitive function, and overall neural health.

By Isla Heckman

References:

- [1] A. Trafton, "This is your brain without sleep," MIT News.
- [2] "Why Do Humans Dream? What Neuroscience Reveals," ScienceInsights, Mar. 18, 2026.
- [3] H. Patel, "The Pons," Teachmeanatomy.info, 2018.
- [4] M. Bellesi, "Sleep and Oligodendrocyte Functions," vol. 1, no. 1, pp. 20–26, Jan. 2015.
- [5] D. Pacheco and A. Singh, "What Happens during NREM Sleep?," Sleep Foundation, Jan. 12, 2022.
- [6] National Institute of Neurological Disorders and Stroke, "Brain basics: Understanding sleep," Sep. 05, 2024.
- [7] Wikipedia Contributors, "Rapid eye movement sleep," Wikipedia, Mar. 21, 2019.
- [8] E. Suni and A. Dimitriu, "Dreams: Why We Dream & How They Affect Sleep," Sleep Foundation, May 02, 2024.
- [9] "What Happens in REM Sleep? Brain, Body, and Dreams," ScienceInsights, May 8, 2026.
- [10] Mark F. Bear, Barry W. Connors, Michael A. Paradiso "Rhythms of the Brain," in Neuroscience: Exploring The Brain (Second edition), ed. Baltimore, Maryland, USA: Lippincott Williams & Wilkins, 2001, ch. 19, pp. 607–626.



Modeling Behavioral Interaction Features in AI-Driven Learning Systems

Acknowledgement

This work was supported by a KING CAROL I-st research fellowship granted by the National Research Authority, contract number 30382/09.02.2026.

1.0 Introduction

Artificial Intelligence (AI) is fundamentally reshaping education, yet its success hinges on how learners engage with AI tools rather than their availability. Understanding behavioral interaction patterns is critical for evaluating both learning outcomes and long-term system adoption [1].

While prior research emphasizes performance metrics, it often overlooks behavioral features like engagement intensity, persistence, and efficiency. These signals show how students balance effort, AI use, and outcomes during study sessions [2].

This study investigates the role of behavioral metrics in predicting user satisfaction and continued usage. The study detects whether the user employed the AI tool again. Using simulated study data and constructing normalized indicators, for example, engagement, persistence, efficiency, and friction tolerance, makes it possible to evaluate their impact on retention. This shows efficiency as the

dominant predictor of continued usage, suggesting that users prioritize outcome quality over raw interaction volume. These findings provide a framework for designing more effective, user-centered AI learning systems.

2.0 Related work

Current learning analytics often rely on simple metrics such as time-on-task to predict performance. However, these models frequently overlook complex dynamics. This study addresses this research gap by introducing a unified framework that integrates engagement, persistence, efficiency, and friction tolerance into a single predictive model.

3.0 Dataset and Methodology

3.1 Feature Engineering

To ensure comparability across variables with different scales, raw interaction features are standardized using Z-score normalization.

$$Z(x) = \frac{x - \mu}{\sigma}$$

where μ and σ denote the mean and standard deviation of the feature, respectively.

Outcome-related variables are transformed using Min-Max scaling to map values to the interval [0,1].

$$MinMax(x) = \frac{x - x_{min}}{x_{max} - x_{min}}$$

This ensures strictly positive values, which are required for ratio-based feature construction.

For stability in division operations, a small constant $\epsilon=0.01$ is introduced to prevent division by zero. Additionally, the engagement metric is shifted to ensure strictly positive values.

$$Engagement_{shifted} = Engagement + |\min(Engagement)| + \epsilon$$

These transformations ensure numerical stability and prevent scale-related distortions in the computed behavioral metrics.

3.2 Data

The data represents a synthetic dataset that models student interaction with AI-assisted learning systems. The dataset contains 11 fields describing user behaviour during AI-assisted study sessions. The dataset used in this study was obtained from Kaggle [3].

Before any analysis was completed, the data was cleaned meaning unnecessary fields were removed and non-numerical fields were replaced with numerical representations. Behavioral features, including engagement, persistence, efficiency, and friction tolerance were computed, and each feature was constructed to positively or negatively weight behavioral signals [4].

3.3 Experimental Setup

A random forest classifier, a machine learning algorithm which constructs and aggregates decision trees [5], was trained to predict continued usage, binary-return status, using an 80/20 stratified split. Optimized via randomized hyperparameter search and cross-validation, the final model utilized 100 trees with maximum depth 10 and Gini impurity, and a fixed random seed for reproducibility. Performance was assessed using accuracy, Receiver Operating Characteristic - Area Under the Curve (ROC AUC), precision, and recall to account for class imbalance. The full preprocessing and training pipeline is documented in the accompanying Kaggle notebook [3].

This study is limited by its use of synthetic data, predefined feature assumptions, and a single-dataset evaluation, which may affect real-world generalizability. However, the proposed framework offers a robust model for behavioral analysis. Future work should validate these findings with real-world datasets and investigate temporal behavior over longer learning periods.

4.0 Behavioral Model

The behavioral dimensions considered in

this study are conceptually grounded in prior work on the ThinkR platform [6], where interaction-based features were shown to reflect cognitive and motivational aspects of learning behavior.

4.1 Engagement (Interaction Intensity)

Engagement captures the total volume of interaction, balancing both time spent and prompt frequency. The engagement feature rewards each prompt and each minute spent studying. Engagement is defined as:

$$Engagement = Z(TotalPrompts) + Z(SessionLengthMin)$$

4.2 Persistence (Cognitive Effort)

Persistence measures sustained effort; it gently penalizes over-reliance on AI assistance, which requires less independent cognitive effort, and rewards cognitive effort. The coefficient (0.3) was empirically chosen to moderate the influence of AI assistance, ensuring that assistance moderates but does not dominate the persistence signal. This reflects the assumption that moderate AI use supports learning, while excessive reliance reduces independent cognitive effort. Persistence is defined as:

$$Persistence = Z(TotalPrompts) - (0.3 * Z(AI_AssistanceLevel))$$

4.3 Efficiency (Outcome Relative to Effort)

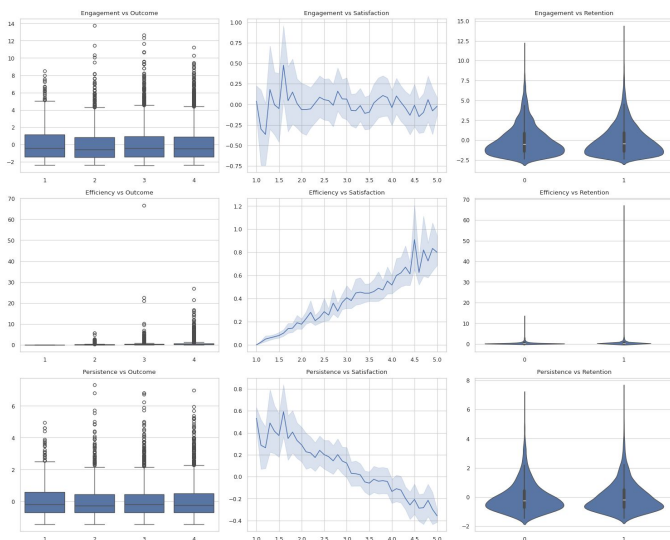
Efficiency computes the successful output relative to engagement; it is modeled as outcome quality relative to interaction cost. Efficiency is defined as:

$$Efficiency = \frac{MinMax(FinalOutcome) * MinMax(SatisfactionRating)}{Engagement_{shifted} + \epsilon}$$

4.4 Friction Tolerance (Pure Effort)

The friction tolerance metric captures cognitive effort by penalizing reliance on AI assistance. Higher values indicate users who engage actively while maintaining lower dependence on automated support. While less expressive than composite features, it serves as a baseline indicator of autonomy in interaction behavior. Friction tolerance is defined as:

$$FrictionTolerance = TotalPrompts - AI_AssistanceLevel$$



These metrics form a structured behavioral representation of user interaction, balancing effort, outcome, and assistance. This feature space enables the subsequent analysis of their relative importance in predicting satisfaction and continued usage.

5.0 Results and Discussion

Following the extraction and normalization of behavioral features, their relationships with learning outcomes, user satisfaction, measured using a 1-5 overall session satisfaction rating, and continued usage were analyzed. Both exploratory analysis and predictive modeling were employed to identify the most influential behavioral signals.

5.1 Behavioral Trends

5.1.1 Engagement

Engagement exhibits a strong positive relationship with both learning outcomes and user satisfaction. Higher engagement scores, derived from increased interaction frequency and session duration, consistently correspond to improved outcomes. This trend suggests that active and sustained interaction with AI-assisted systems contributes positively to task completion and perceived effectiveness. The presence of outliers, captured through Z-score normalization, reflects meaningful behavioral variation rather than noise, indicating distinct user interaction patterns.

5.1.2 Persistence

Persistence demonstrates a clear positive

association with successful task completion. Users who achieve higher outcomes tend to exhibit higher persistence scores, reflecting sustained effort and continued interaction. However, an inverse relationship is observed between persistence and satisfaction. Users reporting the highest satisfaction levels tend to require less persistence, suggesting that prolonged effort may be associated with friction or difficulty during interaction. This highlights a trade-off between effort and user experience: persistence supports success, but excessive persistence may reduce perceived usability.

5.1.3 Efficiency

Efficiency presents the most distinctive behavioral pattern. While lower and intermediate outcome categories show relatively low and stable efficiency values, a sharp increase is observed for sessions resulting in successful task completion. Similarly, higher satisfaction levels strongly correlate with increased efficiency. This indicates that users derive the greatest satisfaction when achieving high-quality outcomes with minimal interaction cost. Efficiency captures the balance between effort and outcome, making it a critical indicator of effective system use.

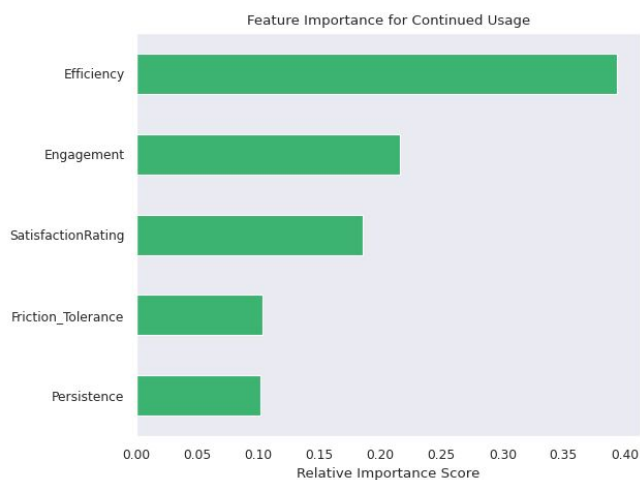
5.2 Predictive Modeling of Continued Usage

The optimized model achieved an accuracy of 72.5% and an ROC AUC score of 0.65, indicating moderate predictive capability. While overall performance is satisfactory, class-level analysis reveals an imbalance in predictive sensitivity. The model achieved a recall of 0.25 for non-returning users, indicating limited detection of disengagement cases. However, precision for this class reached 0.57, suggesting that flagged at-risk users are identified with relatively high reliability.

5.3 Feature Importance Analysis

Feature importance analysis derived from the trained random forest model highlights the relative contribution of each behavioral metric. Efficiency is the dominant predictor, contributing approximately 40% of the model's decision weight. Engagement (22%)

and Satisfaction Rating (18%) serve as strong secondary predictors. Persistence and Friction Tolerance contribute the least (approximately 10% each). These results indicate that continued usage is primarily driven by the efficiency of interaction rather than raw effort. Users are significantly more likely to return when their interactions yield high-quality outcomes relative to the time and effort invested.



6.0 Conclusion

This study proposed a structured framework

for modeling behavioral interaction features in AI-assisted learning systems and evaluated their relationship with user satisfaction and continued usage.

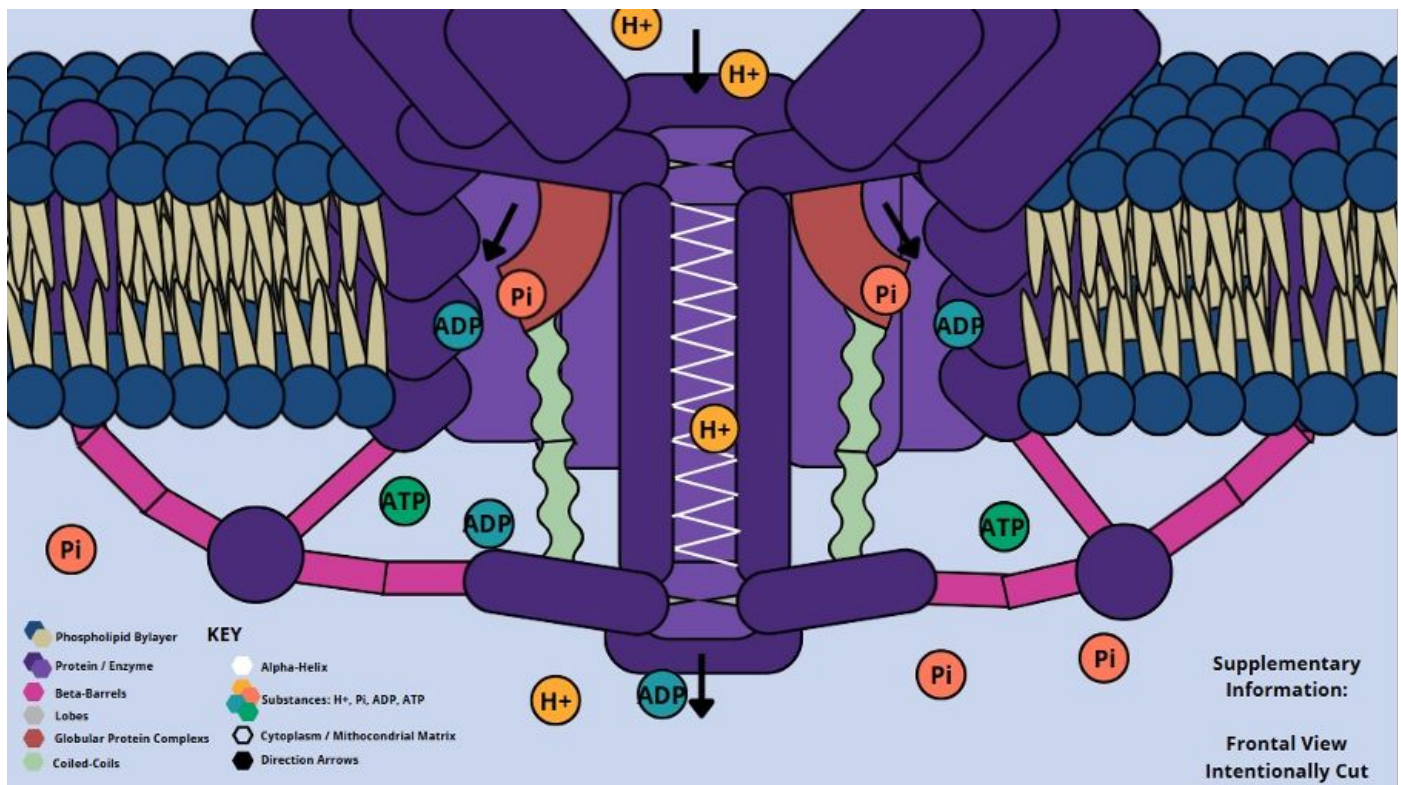
The results indicate that behavioral efficiency is the most significant predictor of system retention, outperforming both engagement and effort-based metrics. This suggests that users prioritize achieving high-quality outcomes with minimal interaction cost rather than maximizing interaction intensity. Accordingly, efficiency emerges as a unifying behavioral signal that captures both performance effectiveness and perceived user experience.

These findings have direct implications for the design of AI-driven educational systems. In particular, systems should prioritize optimizing outcome efficiency and reducing unnecessary interaction overhead, rather than focusing solely on increasing engagement or interaction volume. Such design choices are more likely to support sustained user adoption and perceived usability.

By Iulia Georgiana Buhai

References:

- [1] G. Siemens and R. S. Baker, "Learning analytics: Envisioning a research discipline and a domain of practice," *Proceedings of the 2nd International Conference on Learning Analytics and Knowledge*, 2012.
- [2] J. A. Fredricks, P. C. Blumenfeld, and A. H. Paris, "School engagement: Potential of the concept, state of the evidence," *Review of Educational Research*, vol. 74, no. 1, pp. 59–109, 2004.
- [3] A. Saleem, "AI assistant usage in student life," 2026
- [4] R. S. Baker and P. S. Inventado, "Educational data mining and learning analytics," pp. 61–75, 2014.
- [5] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [6] B. I. Georgiana, "Learning behavior analysis using digital tools," *Penrose Magazine*, vol. 6, no. Issue 6: Neuroscience Mini Issue, pp. 11–14, 2026.
- [7] B. I. Georgiana, "Behavioral Predictors of Retention in AI-Assisted."



Chemiosmosis and Energy Production: Rotation Is Not the Only Way

Premise

Given this is highly speculative work, there are not any specific citations other than generic biochemistry books or articles on the ATP-Synthase, biology, chemistry and physics notions [1], [2].

Introduction

Every living organism requires energy to perform its anabolic and catabolic functions, which allow for creation and destruction respectively, enclosed in what is known as cellular metabolism. This energy, to not cause damage, needs to be stored, commonly in the form of ATP, which is common to every organism. However, while ATP is found everywhere life is, there are only a few production methods found in nature: production at the substrate-level during Glycolysis and Krebs Cycle, production during fermentations to reoxidise crucial coenzymes, and production during cellular respiration.

ATP Production

While production at the substrate-level is often the only way by which an organism can generate energy to survive in the case

of anaerobiosis, cellular respiration is the most efficient one. Cellular respiration utilizes a phenomenon called chemiosmosis [3], where an ionic gradient is created by hydrogen ions, via the electron transport chain (ETC), that then needs to be dissipated. Ions tend to reach an equilibrium when they try to dissipate the gradient created, necessarily passing through the ATP synthase, producing mechanical energy then converted into ATP.

The ATP synthase is a crucial enzyme in this process, without which ATP could not be produced. This enzyme works by a rotational mechanism that converts angular momentum, created by the ions passing through it, in the energy necessary to create ATP. Every ATP synthase discovered to date exploits this rotational process. For example, this process allows for partial conservation of momentum that makes the enzyme spin more easily as time progresses, spending less energy from the ionic gradient; this being one of the reasons why it is omnipresent in nature. Nevertheless, there may be other forms with which ATP can be produced starting from an ionic gradient.

Non-Rotational ATP-Synthase (NRAS)

Starting conceptually from the human ATP

synthase, this enzyme should be theoretically able to produce ATP during cellular respiration without using the rotational movement that is otherwise found in the subunit F₀ of the human enzyme, the transmembrane portion responsible for the rotation.

Contrary to the human enzyme, this NRAS should be a large cylindrical transmembrane complex, partially hollow on the inside. The upper portion of this enzyme, the one that is in contact with the intermembrane space where the ionic gradient is present, contains a diametrically smaller circumference in the center that acts as the “door” by which the ions will pass through via diffusion.

This opening, for mathematical simplicity, possesses three “lobes”, that divide up the circumference, taking direct inspiration from the heart’s tricuspid valve. Ideally, the ions should open those lobes, passing through the circumference, and triggering a chain reaction inside the enzyme. Connected covalently to the lobes on the circumference internal edge, there are globular protein complexes, one per lobe, that present an active binding site for inorganic phosphate (Pi).

On the opposite side of the phospholipid bilayer, the NRAS should have another inner circumference facing the mitochondrial matrix. This second circumference has the same number of lobes as the other, which are to be connected internally to the globular proteins via coiled-coil proteins [4]. These are elastic enough to be able to deform slightly and go back to their standard configurations. Specifically, those coiled-coils should be connected from the lower half of the globular complex proteins near the center of every lower lobe, between the outer edge of the lobe and the center of the overall circumference.

This lower circumference is embedded into a bigger cylindrical structure that “floats” into the mitochondrial matrix. This second cylindrical structure is connected to the main enzyme via the coiled-coils that are connected to the lobes, and is also

anchored to the phospholipid bilayer by other protein structures. These are transmembrane complexes (number-variable) that anchor with β -barrels or similar structures [5], this second cylindrical complex. Furthermore, at a certain point in this polypeptide there should be a smaller globular complex acting as linkage, enabling another chain of β -barrels to generate from it, connecting the main cylindrical transmembrane structure. More NRAS could be connected together to generate supercomplexes (respirasomes).

Both openings are topped by a lid that leaves only a small circle as the lobe’s single point of entrance. The entrance lid is connected to the cylindrical transmembrane protein and also via the lipidic (or amino acidic) insulating strip that divides every lobe, isolating the lobes, reducing proton leak, and permitting asynchronous but parallel opening. The small circle on the center of the circumference, positioned on the tip of the lobes, conveys a more concentrated ion flux, otherwise weaker because of the larger contact area.

The lower circumference’s lid works the same but has n openings, one for each lobe, plus 1. The n openings permit bondage with the coiled-coils complexes to themselves, where the last one lets the ions through. This configuration makes the NRAS unidirectional: reversing it like its rotational counterpart requires some NRASes in their population to be positioned upside-down, adding a protonic pump, like Na⁺/K⁺ ATPase, to the upper circumference’s hole [6]. This may allow some ATP to be produced as well.

ATP Production Process and Mechanistic Considerations

The production process of ATP operated by this specific enzyme can be divided into three phases: Passage, Synthesizing Casual Movement, and Closure and Recharge

In the Passage phase, the ions should stochastically bombard the lobes, pushing onto them, thus opening them and passing through. Whenever a certain number of ions

manage to open the lobes, thanks to the lever principle [7], the lobes themselves, when opening up, push onto the globular proteins that get shoved linearly against the internal wall of the overall cylindrical shape.

Meanwhile, the coils connected to the globular proteins get pulled by this movement, pulling the inferior circumference lobes upward, virtually shutting the enzyme in preparation for the third phase. Additionally, the lever principle amplifies the force applied by the ions necessary to open the lobes. The bigger the lobe itself, the bigger the amplification effect of the lever.

The Synthesizing Causal Movement phase allows the linear displacement of the globular complex that has been pushed by the lobe opening to encounter the inner wall of the overall cylindrical shape of the enzyme. The specific spot contains an active binding site for ADP, that will react with the Pi present onto the globular complex. This force with which the globular complex will be pushed into the walls, should theoretically be sufficient to produce ATP. In terms of scalability, more active binding sites can be positioned in a rack fashion to reach an optimum in production.

This newly-synthesized ATP would now have less affinity to remain bonded into the active binding site of the enzyme's wall, thus being released from it, freeing the active binding site for a new bond. ADP and Pi will enter the enzyme, from the cytoplasmic pool, repositioning themselves in preparation for the next production event.

Finally, the Closure and Recharge phase will take place. In order to completely exit the enzyme, ions have to go through another circumference, this time on the opposite side of the enzyme. This circumference is connected to the globular protein by a coiled-coil complex: from the lower half of the globular complex to the center of the lobes of the lower circumference. Ideally both circumferences, the entrance and the exit, should be synchronized; when one is open, the other one is closed, and vice-versa.

When the ions push onto the lower circumference's lobes, those pull the coiled-coils that in-turn pull onto the globular proteins that get moved back from the ADP's active binding site, "reloading" for the next production. Meanwhile, this causes the upper circumference to close, whilst the lower one gets opened, releasing the ions.

Addendum: Operational Details and Suppletive Considerations

Slowing the ions exiting the NRAS is crucial to prevent premature resetting of the globular complexes, which would result in an unsuccessful synthesizing event. There should thus be present an α -helix connecting both circumferences, which contains polar amino-acid residues inside a cylinder passageway oppositely charged, creating a path that slows ion passage.

Additionally, the lobes are connected to the circumference in two ways: H-bonds and a Polypeptide Pin [8]. The pin passes-through the lobe's thickness, connecting it to the insulated lateral walls; while the H-bonds connect the circumference's edge to the one of the lobe. The H-bonds, number-variable for modularity, get broken by the ion's passage, permitting the rotation of the lobe on the pin axis, letting the ions through.

Further considerations exist, such as the inner part of the upper circumference should have a reversed lid in order to create separate chambers containing the lobes, enabling them to move freely.

Conclusion

This model should, theoretically, be a substitute to rotation although more complicated and difficult to evolve. If proven right or even doable, this project may pave the way for unforeseen biotechnology applications. It could be a way to produce more ATP than the canonical enzyme in the same space inside the inner mitochondrial membrane.

By Jacopo Medici

References:

- [1] T. A. Brow, "Il Ciclo di Krebs e la Catena Respiratoria," in *Biochimica*, Bologna, Italy: Zanichelli (in Italian) 2026, ch. 9, sec. 2, pp. 163–166.
- [2] G. J. Tortora and B. Derrickson, "L'Apparato Cardiovascolare," in *Conosciamo il Corpo Umano - Edizione Azzurra*, 2nd ed. Bologna, Italy: Zanichelli (in Italian) 2020, ch. 8, sec. 5, pp. 239–243.
- [3] "Chemiosmosis," Wikipedia.
- [4] "Coiled coil," Wikipedia.
- [5] "Beta barrel," Wikipedia.
- [6] "Sodium-potassium pump," Wikipedia.
- [7] "Lever," Wikipedia.
- [8] "Hydrogen bond," Wikipedia.



Underwater Drone Engineering

Recently drone engineering has been significantly developed. New technological innovation has relocated the focus of drones into the oceans, allowing for the consistent monitoring of previously endangered underwater ecosystems. This does not include the extra possibilities for transportation, and usage as extensive ocean security. Using different forms of fuel, specifically Hydrogen power [1], long-distance travel of submarine-adjacent drones has become feasible, increasing the efficiency of future environmental consciousness.

In an alarmingly short time period, the global risk to lose 90% or more of our coral reefs by 2050 has risen [2]. Additionally, marine engineering has been unable to properly access and sample the mesophotic layer of reefs on a larger scale for the past 50-60 years. These central tiers are located in an incredibly inconvenient location, as they are too prickly to navigate around but simultaneously too deep to safely send singular divers. Specialized technology for this exists, but is considerably expensive. This issue was overcome when Canadian scholars launched a highly successful drone [1].

Using environmentally conscious hydrogen fuel cells, this vehicle traveled an unprecedented distance of over 1,200 miles

[3]. This triumph met all anticipated expectations, and despite being realistically programmed, surpassed to be able to cover more surface area than ever considered. The layout of this advanced carrier remained extremely lightweight, promising versatility for its future trips. The launch strived for optimal realism, coding the drone for the navigational roadbumps that could occur.

A standard remote operated underwater vehicle (ROUV) must undertake difficulties with pressure, line of sight, and water infiltration through its physical infrastructure, not including malfunctions that can occur during the missions themselves [4]. Combined, these conditions have heavily contributed to stalling in the engineering industry. Additionally, these bodies lack storage capacity for most living beings, and require constant maneuvering and environmentally detrimental fueling sources. However, combined with the new Canadian experimentation and higher demand for environmental solutions, expansion surrounding this topic could provide the assistance Marine Biology is searching for.

Underwater drones are faced with similar issues, but have recently faced different alterations to function. Another type of ROUVs, known as unmanned underwater vehicles (UUVs), became a main focus of the United States Navy in the mid-2010's,

and have since maintained an up-and-coming military development for the nation [5]. Combinations of military determination and professional environmentalism have not always seemed a likely pair. However, the modern world's issues require novel alliances.

Another use for these drones is the commercialization of cross-ocean transportation. The pollutionary expectations of airborne travels are set to triple within the next 25 years [6]. In the hypothetical case that submarine drone programs could be applied to reduce excess CO₂ emissions from planes, this system could be adapted by the shipping and port industry.

Furthermore, the adoption of hydrogen fueling technologies for underwater

transports emerges as a more sustainable new source of energy. As the Envoy submarine drone made its voyage, it created new electricity from the hydrogen fuel while only releasing water as a runoff product [1].

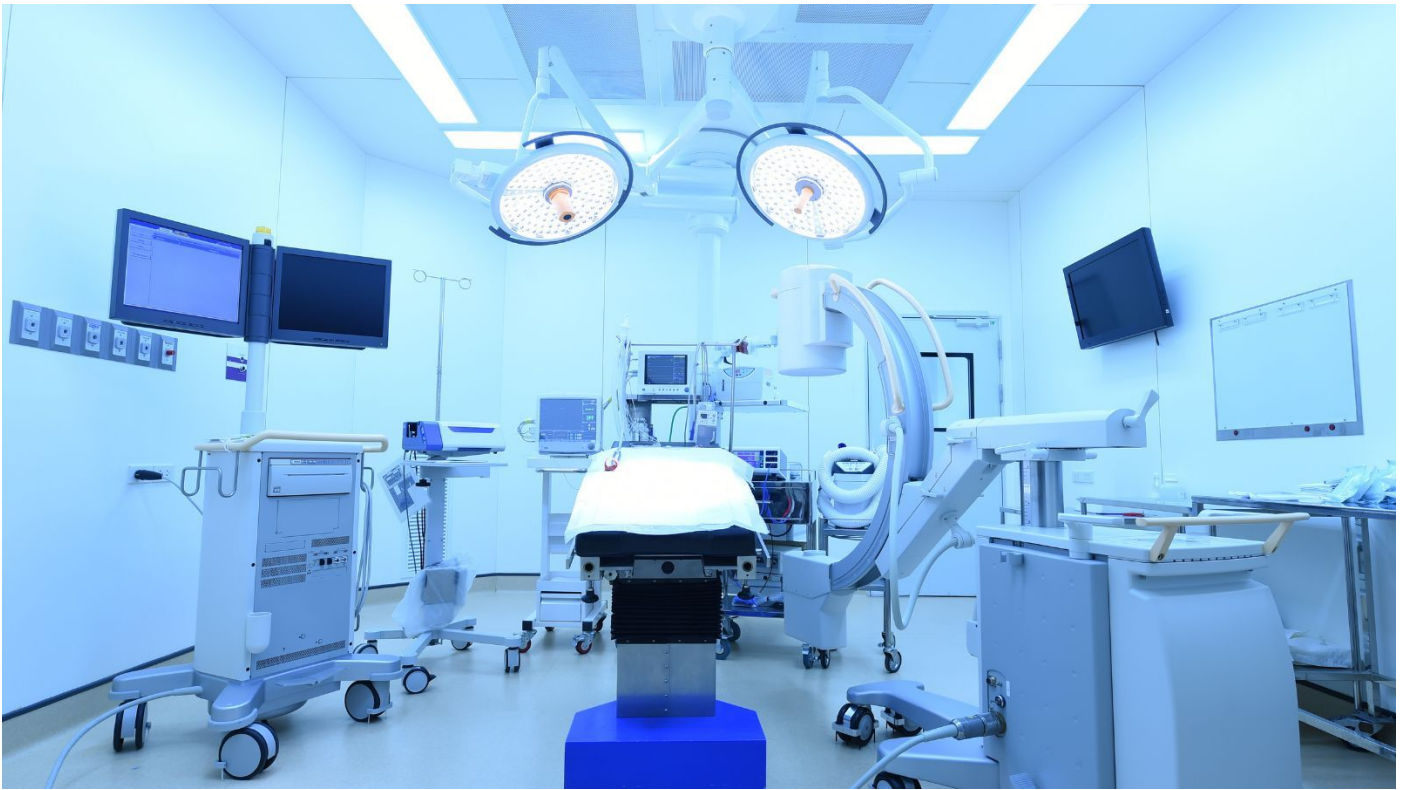
Ocean conservationists, fighting a constant battle against fossil fuels, have voiced support for robotic involvement in marine science work, especially given the capability of hydrogen-powered drones to cleanly produce electricity.

The potential of underwater drones stays promising. Finding proper solutions takes trial, but the extensive research surrounding submarine testing suggests an increase in this combined branch of robotics and environmental science will be achievable.

By Katherine Hatcher

References:

- [1] G. Jedikovska, "[World-first: Submarine drone travels 1,257-miles on hydrogen power](#)," *Interesting Engineering*, Apr. 27, 2026.
- [2] O. Lai, "[5 Coral Reefs That Are Currently Under Threat And Dying](#)," *Earth.org - Past | Present | Future*, Jul. 08, 2021.
- [3] "[How Green Hydrogen Can Keep the Ocean Blue](#)," *Ocean Conservancy*, Sep. 12, 2025.
- [4] "[Remotely operated underwater vehicle](#)," *Wikipedia*, Jun. 13, 2020.
- [5] D. Gettinger, "[What You Need to Know About Underwater Drones](#)," Nov. 16, 2015.
- [6] Wikipedia Contributors, "[Environmental impact of aviation](#)," *Wikipedia*, Mar. 09, 2019.



The Rise Of AI In The Operating Theatre

The nature of operating theatres has changed significantly with modern developments. Today, even the most skilled surgeons may work alongside robotic systems guided partly by artificial intelligence rather than solely by human instinct. Surgical robots have gradually become part of hospitals worldwide, allowing for more precision and control in high-risk operations.

Among the first robotic systems to achieve global recognition was the da Vinci Surgical System, which has now assisted in more than 12 million operations worldwide [1]. During an operation, the surgeon sits at a console and controls several robotic arms equipped with tiny surgical instruments and a high-definition 3D camera. Importantly, the robot does not act independently. Instead, it translates the surgeon's hand movements into smaller, steadier, and more precise actions inside the patient's body.

This system allows surgeons to perform minimally invasive procedures through very small incisions, reducing blood loss, pain, scarring, and recovery times for patients. It also improves precision, especially in delicate operations, because surgeons are provided with a much clearer and more

accurate view of the surgical area.

Traditional robots follow pre-programmed instructions and fixed commands, whereas AI systems are capable of learning through machine learning. In surgery, this allows AI to analyse thousands of previous operations and identify patterns linked to successful outcomes or potential complications.

An example of these modern robotic systems is the CMR Surgical's Versius and Medtronic's Hugo RAS, which combines robotic engineering with artificial intelligence to improve surgical precision [2]. Their algorithms analyse data from previous procedures to refine the stability and accuracy of surgical movements. Their built-in cameras do more than simply display an image of the operation; through computer vision, they can interpret the surgical scene by recognising tissue boundaries, tracking movement, and identifying important structures such as blood vessels or nerves.

This is made possible by computer vision, a branch of AI that processes and interprets visual information. When the robot's endoscopic camera, a flexible tube fitted with a camera and light used to look inside the body, streams a live feed, AI models can identify anatomical structures. This allows

danger zones to be highlighted, and tissue movement can be predicted. The surgeon stays in control, but with a second analytical system helping with every movement.

Some systems extend these capabilities even further. Experimental platforms such as the Smart Tissue Autonomous Robot (STAR), developed at Johns Hopkins University in America, can already perform soft-tissue suturing with minimal human input [3]. The robot does this by analysing tissue elasticity, calculates the best stitch placement, and adjusting its movements accordingly. In one trial, STAR outperformed experienced surgeons in precision and consistency.

For patients, the benefits are tangible. AI-assisted procedures result in smaller incisions, less blood loss, and shorter hospital stays, becoming no longer a hypothetical concept but very measurable clinical outcomes. A systematic review published in *Surgical Endoscopy* found that robotic surgery was associated with lower blood loss, shorter hospital stays, and reduced complication rates compared with open surgery [4].

For example, in complex neurosurgery, AI-enhanced navigation systems can display detailed maps of the brain, marking critical areas to avoid. Furthermore, cardiothoracic surgery algorithms can calculate blood-flow dynamics during procedures.

AI systems are also being developed to analyse surgical sounds. Acoustic analysis algorithms are currently being trained to detect subtle changes in the pitch and frequency of surgical sounds, such as the tone of cutting instruments or suction devices. These changes may indicate complications including internal bleeding, tissue damage, or equipment malfunctions before the human ear is able to notice them.

One of the most significant advantages of AI in surgery is data integration. During an operation, surgeons may need to consult imaging scans, patient histories, and

physiological monitors simultaneously. AI systems use machine-learning algorithms to rapidly compare and combine this information within a matter of seconds, highlighting important patterns or sudden changes that matter most.

However, as artificial intelligence becomes more and more involved in surgical decision-making, concerns surrounding accountability, patient trust, and medical ethics emerge. If an AI-assisted system makes an error during an operation, the allocation of responsibility becomes less clear. Therefore the rise of robotic surgery challenges both the limits of technology and the ethical boundaries that support modern medicine.

Furthermore, training future surgeons presents another challenge. Surgeons of the future may need to master both anatomy and algorithms. As a result, many leading universities now include robotics and data science in surgical training, yet medical schools still often assess students through traditional paper exams and laboratory work rather than coding or sensor-based technologies. Therefore, medical education must continue to evolve alongside advances in artificial intelligence.

There is also the issue of equity. A fully robotic operating theatre can cost millions. Wealthy hospitals in major global cities may offer the latest robotic systems, while smaller regional hospitals and developing countries often struggle to afford even basic laparoscopic tools, instruments used for minimally invasive surgery through small incisions. For example, Sierra Leone in West Africa continues to face major shortages in surgical equipment and healthcare funding [5]. This could result in a new form of “digital divide,” not only between the community, but also between patients who have access to AI-enhanced healthcare and those who do not.

Patient consent also remains a central concern. Patient confidence may be influenced by the knowledge that their surgeon is assisted by an AI system that

continues to improve through data and experience. A recent survey found that while many patients support doctors using AI as a supportive tool, far fewer are comfortable with fully autonomous robotic decision-making in surgery [6]. Human reassurance, empathy, and trust remain essential parts of medicine that no algorithm can fully replicate.

While current systems mostly remain under human control, researchers are beginning to test autonomous microsurgery—procedures performed on extremely delicate areas such as blood vessels, nerves, or the retina of the eye. In ophthalmology, robots can already inject drugs into the retina with sub-millimetre precision. In orthopaedics, AI systems can assist with bone alignment more accurately than traditional manual methods.

The logical endpoint of this technology is “closed-loop surgery,” where an AI system continuously monitors patient data, adapts its actions during the procedure, and carries

out parts of surgery with minimal human intervention. Although this concept may appear futuristic, it is based on the same machine-learning principles that are already used today, such as in self-driving cars and automated robotics.

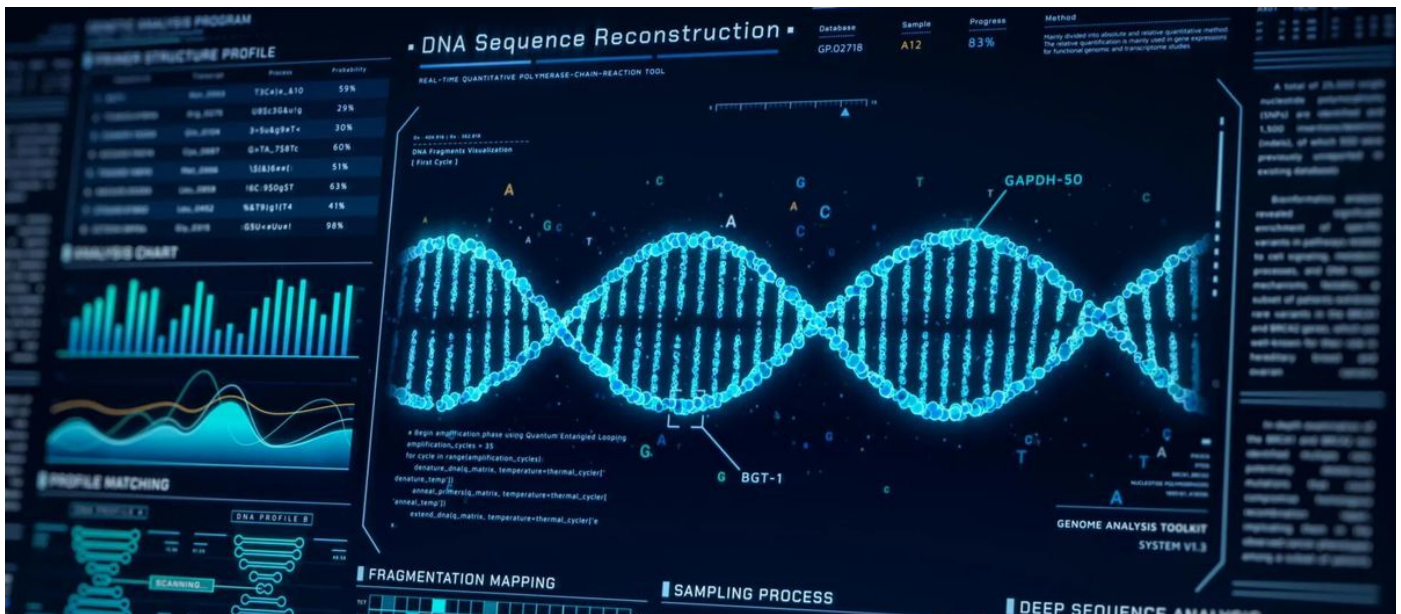
Interestingly, AI does not necessarily make surgery colder or less human. Some systems use augmented reality, a technology that overlays digital images and data onto the surgeon’s field of view during an operation. This can highlight where to cut, identify hidden blood vessels or nerves, and improve surgical accuracy. In many ways, it allows surgeons to visualise anatomy more clearly than ever before.

Ultimately, AI does not replace the surgeon, but supports them. While surgeons contribute judgement, experience, and communication with patients, AI can improve precision and maintain consistent performance during long and complex procedures.

By Kiana Oliver

References:

- [1] Intuitive Surgical, "[da Vinci Surgery](#)," Intuitive Surgical.
- [2] CMR Surgical, "[Versius Surgical Robotic System](#)," CMR Surgical.
- [3] J. Shademan et al., "Supervised autonomous robotic soft tissue surgery," *Science Translational Medicine*, vol. 8, no. 337, 2016.
- [4] A. Author et al., "Robotic versus open surgery outcomes: systematic review," *Surgical Endoscopy*, vol. 31, no. 2, pp. 1-15, 2017.
- [5] World Health Organization, "[Global surgery and healthcare access](#)," WHO.
- [6] National Library of Medicine, "[Public attitudes toward AI-assisted healthcare and surgery](#)," PubMed.



Predicting the Invisible: AI in Non-Coding DNA Research

Introduction

For much of the 20th century, molecular biology was grounded in a simplified conceptual model in which genes encode proteins, and proteins determine phenotypic traits [1]. It has shaped how scientists understood inheritance and disease. If a heritable disorder existed, determining the cause to be a faulty gene was a rational conclusion. However, as genomics advanced and the Human Genome Project (HGP) was developed, this idea was challenged. Approximately 2% of the human genome codes for protein, while the majority of the genome is non-coding [2]. This highlighted a fundamental gap in understanding the role of the vast non-coding genomic regions.

Central Dogma

The central dogma is the framework guiding molecular biology: DNA is replicated, transcribed into RNA, and translated into proteins. Proteins are responsible for the majority of cellular functions and determine observable traits, making this framework successful in explaining inherited diseases and metabolic pathways [3].

Advances in genome sequencing, completion of the HGP, and progress in Genome-Wide Association Studies (GWAS) gradually challenged this protein-centric

perspective. These projects revealed that most disease-associated variants lie in non-coding regions of DNA [4], [5].

These variants, located far from the genes they influence, are tissue-specific, context-dependent, and difficult to interpret. Unlike coding mutations, non-coding variants do not directly alter protein structure [5].

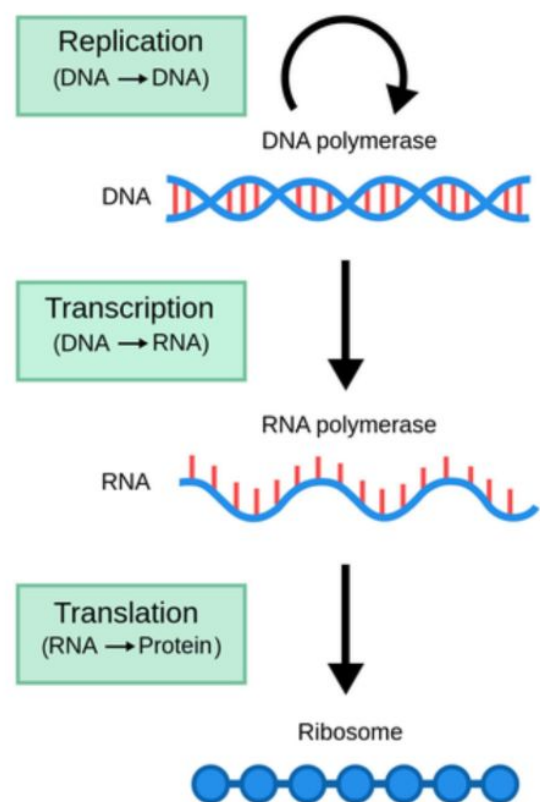


Figure 1: Central Dogma of Molecular Biology [6].

Functions of Non-Coding DNA

Initially dismissed, non-coding DNA is now

understood to carry out diverse biological functions. Promoters, enhancers, and silencers regulate when and where genes get activated or deactivated. Functional RNAs like tRNA, mRNA, microRNA, lncRNA, and eRNA are produced. Certain base pair sequences serve as origins of DNA replication or form Untranslated Regions (UTRs) at the end of mRNAs and regulate protein functions [7].

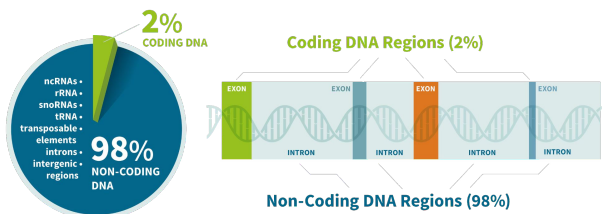


Figure 2: Illustration of non-coding DNA [8].

Importantly, regulatory elements influence genes located thousands or hundreds of thousands of base pairs away from them. Non-coding DNA's functions depend not only on sequence but also on cell type, developmental stage, chromatin structure, and environmental context [7]. Traditional experimental approaches are expensive, slow, and limited in scale. This is where artificial intelligence has begun to transform the field.

AI in Play

AI has become a powerful tool for predicting the functions of non-coding DNA. These models enable *in silico* mutagenesis—a computational method that simulates genetic mutations to predict their functional effects [9].

Advances in deep learning have led to increasingly sophisticated models, notably DeepSEA, Basenji, Enformer, and AlphaGenome. These multimodal AI models are capable of processing and integrating information from multiple modalities, thereby enabling the prediction of genomic features [10].

According to [11], Convolutional Neural Networks (CNNs) are effective at filtering and processing large datasets. They exhibit superior pattern-recognition skills. In this context, they help identify the cause-effect

relationships arising from changes in the DNA base-pair patterns. CNNs generally consist of two primary types of layers: convolutional layers, which process and transform input information; and non-convolutional layers, which assist convolutional layers. Pooling techniques such as Max pooling, Average pooling, and Attention pooling are used to balance resolution with computational efficiency [12], [13].

DeepSEA (2015)

DeepSEA, a deep learning algorithmic framework, processes 1000 base-pair sequences and focuses on local regulatory signals. By predicting how single-nucleotide changes affect chromatin features, DeepSEA helps identify genetic variants that are likely to have functional effects.

DeepSEA first takes a 1 kb input and transforms it using its CNN architecture with three main layers: Convolutional Layers acting as scanners looking for motifs, Pooling Layers (Max pooling), and Fully Connected Layers integrating patterns to make a final prediction. Trained on ENCODE and Roadmap Epigenomics datasets, DeepSEA demonstrated *de novo* predictions—predictions based strictly on the DNA sequence without the need for prior variant knowledge—supporting the analysis of rare mutations and *in silico* mutagenesis [14].

Limitations of DeepSEA

DeepSEA's short input window prevents it from capturing long-range interactions. It also treats DNA as a linear sequence, ignoring 3D chromatin organization and Topologically Associating Domains (TADs). Like many CNNs, DeepSEA involves complex mathematical computations. Its binary outputs limit interpretability, and training biases toward well-studied cell types reduce generalizability [15].

Basenji (2018)

Basenji, successor to the “Basset” model, accepts input sequences up to 130 kilobases, allowing it to capture distal

interactions. To overcome the massive data output that would typically result from this input size, Basenji uses dilated convolutions. It first scans the DNA at 1 bp resolution to identify small motifs, then applies a max-pooling process to compress the sequence into 128 bp bins. After sequence compressions, Basenji applies dilated convolutions. The resulting high-resolution map of regulatory activity covers a massive genomic distance.

Basenji, trained simultaneously on thousands of datasets, as well as different cell types and species, is more accurate. Basenji uses a Poisson likelihood model, predicting a continuous signal track, allowing researchers to quantify the magnitude of a mutation's impact, rather than just stating that an impact exists [16].

Limitations of Basenji

Basenji's accuracy beyond the 20 kb window declines, limiting its ability to capture distal regulatory elements, while its pooling and dilated convolutions reduce effective resolution, limiting precise localization of regulatory elements. Its sensitivity to sequence variations may inflate predicted effects for variants without real phenotypic impact, requiring experimental validation. The model's ability to capture individual-specific variation is limited, as it is primarily trained on reference genome [17].

Enformer (2021)

Enformer, a portmanteau derived from "Enhancer" and "Transformer", combines the local scanning power of traditional biology models and the memory of modern AI models. Enformer takes an input of 196608 bp of DNA sequence.

Enformer's convolutional blocks pool the data to 128 bp, reducing spatial dimension. Transformer blocks then capture long-range interactions across sequences, while a final cropping step trims 320 positions on each side to avoid computing loss at the far ends. The output heads generate predictions for organism-specific tasks.

While both Basenji and Enformer use convolutional blocks, Enformer's Transformer architecture overcomes the limitations of dilated convolutions. Trained on both human and mouse genomes, Enformer is more generalizable and accurately predicts gene regulation for unseen sequences. It performed exceptionally well by predicting the effects of variants over 50 kb away. Using attention pooling, the model prioritizes the genomic regions most relevant to specific traits [17].

Limitations of Enformer

Enformer is computationally expensive and memory-intensive. While it highlights important regulatory regions, it does not clearly explain the underlying biological processes and often shows sign inconsistencies. Enformer also struggles to model environmental influences and lacks fine resolution within the same 128 bp bins [17].

AlphaGenome (2025)

AlphaGenome addresses the challenge of handling long input sequences while preserving resolution by using a U-Net-style architecture that processes up to one million base-pairs—an advance over earlier models.

Beginning with the Encoder, AlphaGenome progressively downsamples the input from 1 bp to 128 bp resolution, while simultaneously increasing the number of channels to enhance representational capacity. The data is then processed in the Transformer Tower, where Sequence Parallelism divides the sequence into smaller segments, with each segment processed by a different TPU core.

Within this stage, Pairwise Interaction Blocks (PIBs) update a pairwise representation that captures interactions between genomic positions. The information generated at PIBs is sent back to the tower in the form of a bias signal, encouraging the model to focus more strongly on interactions between these physically proximal sites. Contact Maps are also

predicted here. Subsequently, the Sequence Decoder reverses the operations of the Encoder by upsampling the data from 128 bp to 1 bp resolution. Finally, task- and organism-specific Output Heads generate corresponding predictions.

AlphaGenome—trained on massive human and mouse genomes from various sources such as ENCODE, GTEx, 4D Nucleome, and FANTOM5—used only half the computational power of Enformer. Its use of average pooling and attention pooling mechanisms improves accuracy [19].

Limitations of AlphaGenome and their implications

AlphaGenome's accuracy declines for regulatory interactions beyond 100 kb and remains limited to human and mouse genomes. While it predicts molecular effects well, connecting these predictions to complex traits and condition-specific responses remains challenging [19].

Additionally, although AlphaGenome has demonstrated capability in inferring cause-and-effect relationships, these models remain primarily predictive rather than causal. Their performance depends heavily on training data, inheriting biases

and gaps. Accurately modeling gene regulation across diverse cell types, environments, and developmental stages remains a challenge.

Interpreting non-coding DNA could transform medicine by enabling the development of personalized therapies that target gene regulation rather than proteins, which are often accompanied by side effects. It could improve genetic screening and allow more precise diagnosis of inherited diseases. Technologies that can reprogram cells, particularly in plants to improve their resistance to environmental stress, could also be developed.

Conclusion

AI has reshaped how scientists approach the genome. By uncovering patterns hidden within non-coding DNA, AI models have drastically expanded our knowledge of gene regulation and disease. The future of genomics will depend not only on increasingly powerful models, but also on deeper insight into what those models capture and how their predictions relate to biological systems. In uncovering the functions of the genome's non-coding majority, AI has provided biology with new analytical frameworks.

By Lakshmi N Gowda

References:

- [1] F. H. Crick, "On protein synthesis," in *Symposium of the Society for Experimental Biology*, vol. 12, pp. 138–163, 1958.
- [2] F. Moraes and A. Góes, "A decade of human genome project conclusion: Scientific diffusion about our genome knowledge," *Biochemistry and Molecular Biology Education*, vol. 44, no. 3, pp. 215–223, 2016.
- [3] C. L. Tan and E. Anderson, "The new central dogma of molecular biology," *Resonance*, vol. 14, no. 3, pp. 1–32, 2020.
- [4] H. Giral, U. Landmesser, and A. Kratzer, "Into the wild: GWAS exploration of non-coding RNAs," *Frontiers in Cardiovascular Medicine*, vol. 5, p. 181, 2018.
- [5] Y. Zhu, C. Tazearslan, and Y. Suh, "Challenges and progress in interpretation of non-coding genetic variants associated with human disease," *Experimental Biology and Medicine*, vol. 242, no. 13, pp. 1325–1334, 2017.
- [6] Labster, "From DNA to Protein," image.
- [7] R. P. Alexander et al., "Annotating non-coding regions of the genome," *Nature Reviews Genetics*, vol. 11, no. 8, pp. 559–571, 2010.
- [8] AncestryDNA, "Non-coding DNA," image.
- [9] H. Zhang et al., "NCNet: Deep learning network models for predicting function of non-coding DNA," *Frontiers in Genetics*, vol. 10, p. 432, 2019.
- [10] X. Wang, "The evolution of multimodal AI: Creating new possibilities," *Int. J. Artif. Intell. Sci. (IJAI4S)*, vol. 1, no. 2, 2025.
- [11] K. O'Shea and R. Nash, "An introduction to convolutional neural networks," *arXiv, preprint arXiv:1511.08458*, 2015.
- [12] L. Zhao and Z. Zhang, "An improved pooling method for convolutional neural networks," *Scientific Reports*, vol. 14, no. 1, p. 1589, 2024.
- [13] H. Gholamalizadeh and H. Khosravi, "Pooling methods in deep neural networks: A review," *arXiv, preprint arXiv:2009.07485*, 2020.
- [14] J. Zhou and O. G. Troyanskaya, "Predicting effects of noncoding variants with deep learning-based sequence model," *Nature Methods*, vol. 12, no. 10, pp. 931–934, 2015.
- [15] J. Stachowicz, "Exploring DeepSEA CNN and DNABERT for regulatory feature prediction of non-coding DNA," 2021.
- [16] D. R. Kelley, Y. A. Reshef, M. Bileschi, et al., "Sequential regulatory activity prediction across chromosomes with convolutional neural networks," *Genome Research*, vol. 28, no. 5, pp. 739–750, 2018, doi: 10.1101/gr.227819.117.
- [17] Ž. Avsec, V. Agarwal, D. Visentin, et al., "Effective gene expression prediction from sequence by integrating long-range interactions," *Nature Methods*, vol. 18, pp. 1196–1203, 2021, doi: 10.1038/s41592-021-01252-x.
- [18] DeepMind, "Enformer model architecture," image.
- [19] Ž. Avsec et al., "AlphaGenome: advancing regulatory variant effect prediction with a unified DNA sequence model," *bioRxiv*, Jun. 2025.
- [20] DeepMind, "AlphaGenome animation," video.

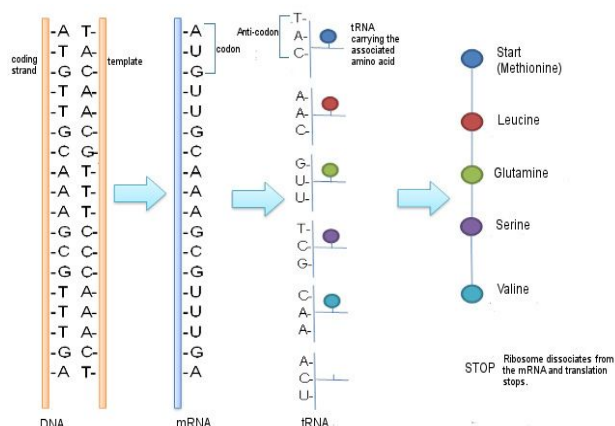


DNA Phenotyping: Applications, Limitations, and Ethical Challenges

Introduction

Creating art using DNA from individuals of the past or present is an intricate process using and analysing genetic data. DNA fingerprinting, discovered by Sir Alec Jeffreys in 1984, was the beginning of the advancements of DNA utilisation in the field of forensics.

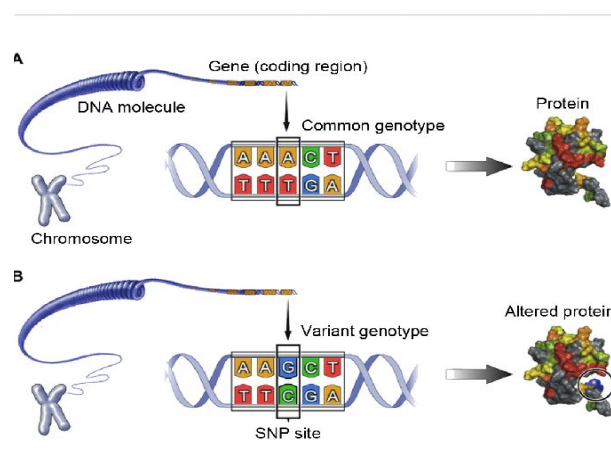
How Phenotypes Are Related To Genotypes



DNA phenotyping, also known as molecular photofitting, is a forensic technique that predicts an individual's externally visible characteristics (EVCs) and biogeographical ancestry from DNA [1]. A phenotype is the observable expression of the genotype, or genome. Genes produce specific phenotypes by acting as blueprints for synthesising proteins and RNA molecules through transcription and translation. These proteins act as enzymes, structural

components or signalling molecules that dictate cellular structure and function; thus resulting in a unique appearance.

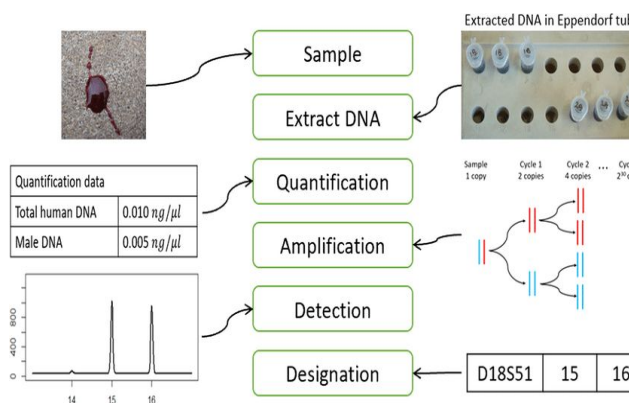
Forensic DNA phenotyping looks at certain regions of the genome that are directly or indirectly associated with normal variation in physical appearance among individuals [1]. The process relies heavily on the utilisation of the existence of single nucleotide polymorphisms (SNPs), which are the most common type of genetic variation among people. An SNP is the result of a single nucleotide (A, T, C, or G) being replaced by another nucleotide [2]. SNPs act as distinct, high resolution genetic markers across the genome and are located using specific loci. SNPs help identify variations associated with diseases, traits or ancestry. They have been officially used in forensics since 2014. However, to be classified as an SNP, a variation must be present in at least one percent of the population [2].



Process of DNA Phenotyping

The methodical process of DNA phenotyping includes six key steps. DNA is first gathered from various sources, such as blood, saliva, bone, hair and teeth. Prior to further analysis, DNA quantity and degradation must be assessed. DNA is then extracted from cells, purified to remove any contaminants and separated from cellular debris. The extracted DNA undergoes a process known as polymerase chain reaction amplification, where a specific segment of DNA is copied to generate sufficient quantities, typically reaching billions of copies.

Within the amplified DNA, SNPs associated with physical traits are identified. Technology known as massively parallel sequencing (MPS) is utilised to analyse hundreds of markers in a single assay [3]. Bioinformatic analysis is then applied to identify the specific SNP genotypes. The genotypes are then run through predictive models to generate probabilities for specific traits. This classifies DNA phenotyping as a probabilistic predictive process.



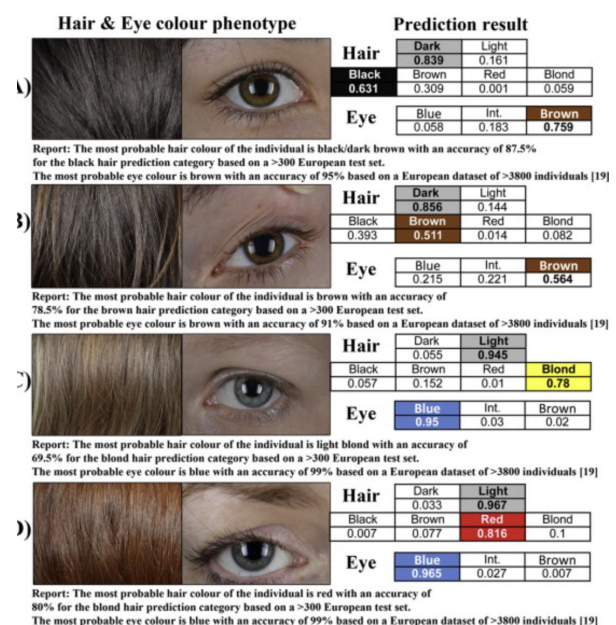
How Dna Phenotyping As A Predictive Process Is Utilised

Using current technology, only a limited number of features can be predicted: variations of pigmentation, partial face morphology, biogeographical ancestry, age estimation, and sex determination. When determining the colours of specific features, scientists will analyse the amount, distribution and type of melanin pigment present. Eumelanin is responsible for darker shades, such as brown eyes, darker hair and skin. Pheomelanin is responsible for

lighter shades including red hair, freckles and fair skin. Pheomelanin provides very little UV protection [1].

When working with pigmentation, predictions range from a 99% accuracy for eye colour, 64%-95% for hair colour and 72%-99% for skin colour [4]. To assist in the prediction of eye colour, researchers have created systems such as the original IrisPlex, which consists of six SNPs present among genes such as (HERC2, OCA2, SLC24A4, SLC45A2/MATP, TYR, IRF4) [1]. The HirisPlex-S system, which enables the simultaneous prediction of human eye, hair, and skin color from DNA, consists of 41 Single Nucleotide Polymorphisms (SNPs) [3].

Age estimation can also be performed using DNA methylation analysis [1]. DNA methylation patterns change predictably with age. These changes can be detected and used to determine an individual's age from biological samples with as few as seven markers and a variety of origins (tissues and bodily fluids), in multiple contexts.



DNA phenotyping cannot be used alone when determining the face morphology of an individual. Companies like Parabon combine DNA predictions with skull morphology and tissue depth markers applied to the decedent's skull [3]. A final composite is produced by combining the

two predictions digitally. Face morphology is polygenic and is influenced by the environment, developmental randomness, age and epigenetics; therefore, the utilisation of DNA alone is not sufficient.

Applications and ethical concerns

DNA phenotyping has been used since the early 2000s in criminal investigations, identifying unknown remains, as well as in archaeological reconstructions. Although this process has proven effective, there are ethical concerns associated with it. DNA phenotyping generates a general overview of the potential characteristics of an individual; hence, there is a risk of misidentification, as the sophisticated prediction is not precise and may resemble a number of individuals.

Phenotyping may also perpetuate societal prejudices by stigmatising specific racial or ethnic groups by focusing on predicted ancestry. In the field of forensics, specifically in obtaining crime scene material; the fugitive is not always present. There is a possible violation of privacy as sensitive information about suspects or victims may be exposed without their consent.

Additionally, the surreptitious collection of DNA without consent poses an ethical concern, especially if the sample is used for purposes other than the initial investigation. The technology used in the DNA phenotyping process may have a degree of bias, as the models are trained on limited datasets and are less accurate for underrepresented populations, leading to complications.

Forensic genetics research has to adhere to

certain protection acts in order to maintain an ethical status. These laws, such as the General Data Protection Regulation and Protection of Freedoms Act 2012, combined with the approval from independent ethical review boards have become standard practice since 2020 [5].

Conclusion



In conclusion, molecular photofitting is an active and established method for predicting the externally visible characteristics of an individual. Combining DNA predictions with skull morphology and tissue depth models is the modern scientific art of generating facial approximations. It is imperative to refrain from complete reliance on DNA phenotyping as a process, as predictions depend on population data, statistical models and an incomplete genetic understanding. With the current advances in AI and technology, there may be potential for improvements in accuracy.

By Madison Krause

References:

- [1] M. A. Rahat et al., "Forensic DNA Phenotyping," Dec. 2, 2022.
- [2] MedlinePlus, "What are single nucleotide polymorphisms (SNPs)?," U.S. National Library of Medicine.
- [3] C. Xavier C et al., "Development and validation of the VISAGE AmpliSeq basic tool to predict appearance and ancestry from DNA," Forensic Science International: Genetics, 2020.
- [5] M. Kayser, "Forensic DNA Phenotyping: Predicting human appearance from crime scene material for investigative purposes," Forensic Science International: Genetics, 2015.



The Ghost in the Machine: Learning to Live with a Prosthetic Limb

Introduction

For centuries, human augmentation was limited to inanimate objects. From the wooden toes used in ancient Egypt to the high-tech carbon-fiber blades used by today's athletes, prosthetics were essentially tools. They were objects controlled by the user, restricted by the laws of physics and the user's conscious effort.

The combination of artificial intelligence (AI) and advanced neural interfaces has created a new category of bionic limbs. These devices do not wait for a command; they actively participate in movement. This allows for "neuro-bionic symbiosis," where the human nervous system and artificial limbs blend to work as a single unit, raising deep questions regarding both the limits of technology and the altering of human identity upon integration.

When Prosthetic Moves First

The reality of an AI-driven prosthetic can be understood through the "feedforward" loop. In a natural limb, the brain sends an electrical signal through the spinal cord to the nerves, which makes a muscle contract. Modern bionics work similarly: sensors placed on the remaining part of the limb read the electrical activity of the muscles in

a process called electromyography (EMG) [1].

This creates a philosophical puzzle regarding agency. If a machine acts for you by anticipating your will, agency becomes divided.

The human mind initiates the intent, while the predictive software executes the physical response. This collaboration makes both the user and the system co-authors of the action.

Beyond Function

Traditionally, even the best prosthetics have been considered extracorporeal—tools attached to the body but not part of it [3]. This changed with closed-loop systems which take data from sensors in the bionic fingertips, like touch and pressure, and send it back to the user's nerves [4].

Restoring the sense of touch is revolutionary. Research supported by the National Institutes of Health (NIH) shows that this feedback remaps the brain's peripersonal space—the mental map of the area immediately surrounding our bodies [5]. Without touch, the brain sees the prosthetic as an object. With touch, the brain expands its definition of "self" to include the machine. This improves the user's body image and reduces the psychological weight of disability much

more effectively than movement alone.

Phantom Feelings and the Algorithm

Many patients describe a sudden realization where the heavy, mechanical device seems to disappear from their mind, replaced by the natural feeling of simply having an arm again [6]. Interestingly, this integration often connects with a patient's "phantom limb." Instead of clashing with the sensation of the missing hand, the AI device telescopes—the brain's internal map stretches to fit the physical dimensions of the new carbon-fiber and titanium fingers.

Identity Rupture vs. Symbiosis

While many welcome the integration of bionics, others feel a sense of self-estrangement. Instead of feeling "whole," they feel like a biological host carrying an alien presence [7]. Some scholars suggest the concept of "cyborg multiplicity"—the idea that these individuals are not half-human and half-machine, but a

new, single type of identity [8].

This situation introduces complex legal and ethical challenges. If a drop of sweat causes a sensor to malfunction and the AI crushes a glass, determining liability between the user and the software is difficult. Current legal frameworks are unprepared for this because they assume a person is a single biological unit [9]. Our current laws, which assume a person is a single biological unit, are not ready for this.

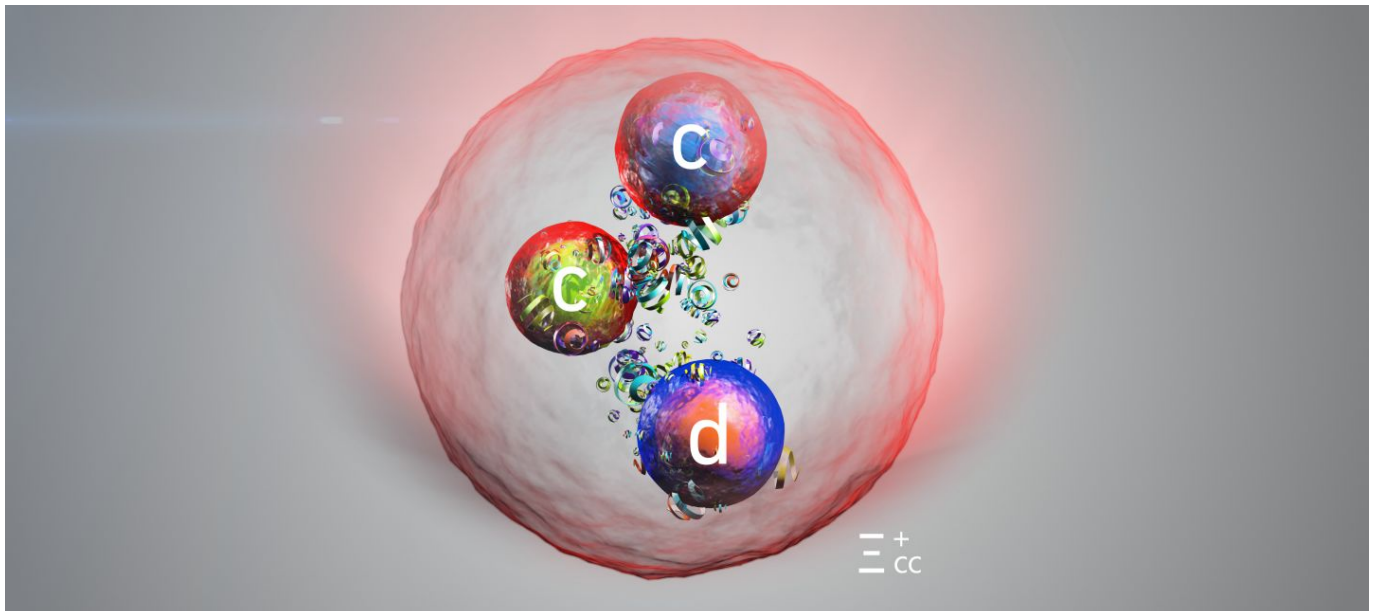
Conclusion

The bionic limb is the ultimate version of the extended mind. AI-driven prosthetics do not diminish humanity; instead, they reveal its inherent flexibility. By bridging the gap between human nerves and predictive algorithms, these tools serve as a visible extension of our biological evolution. Building machines that read thoughts and interact with the physical world does not replace the human self, it simply realizes how expansive it has always been.

By Maimouna Tougoutcho Coulibaly

References:

- [1] T. Kuiken et al., "Targeted muscle reinnervation for real-time myoelectric control of multifunction artificial arms," *JAMA*, vol. 301, no. 6, pp. 619–628, 2009.
- [2] J. M. Hahne et al., "Linear and nonlinear regression techniques for simultaneous and proportional myoelectric control," *IEEE Transactions on Neural Systems and Rehabilitation Engineering*, vol. 22, no. 2, pp. 269–279, 2014.
- [3] M. Botvinick and J. Cohen, "Rubber hands 'feel' touch that eyes see," *Nature*, vol. 391, no. 6669, p. 756, 1998.
- [4] S. Raspopovic et al., "Restoring natural sensory feedback in real-time bidirectional hand prostheses," *Science Translational Medicine*, vol. 6, no. 222, pp. 222ra19, 2014.
- [5] G. Di Pino et al., "Neuroplasticity and embodiment of a bionic hand: A peripersonal space perspective," *Journal of NeuroEngineering and Rehabilitation*, vol. 11, no. 4, 2020.
- [6] P. Marasco et al., "Illusory movement reproduction in amputees with artificial limbs," *Brain*, vol. 140, no. 11, pp. 2838–2848, 2017.
- [7] F. Lucivero and A. Tamburrini, "Ethical implications of neuro-prosthetics: Identity rupture and self-estrangement," *Neuroethics*, vol. 12, pp. 1–14, 2019.
- [8] O. Belton, "Representations of Posthuman Women in Contemporary Science Fiction Television," Ph.D. dissertation, University of East Anglia, 2013.
- [9] A. M. Farrell and S. Quigley, "Liability and fractured agency in AI-driven closed-loop bionics," *Journal of Law and Biosciences*, vol. 8, no. 1, 2021.



Proton, but Heavier: New Subatomic Particle Discovered in The Large Hadron Collider

Introduction

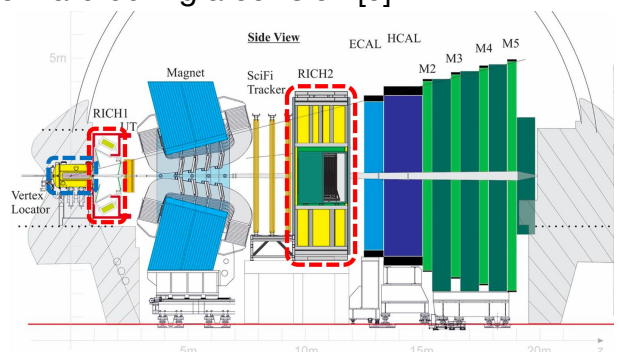
On March 16th, at the Rencontres de Moriond 2026 Electroweak conference, CERN scientists announced the discovery of a subatomic particle. The particle was given the name Xi-cc-plus, also denoted as Ξ_{cc}^+ . By virtue of its composition, Ξ_{cc}^+ is also referred to as the “heavy version of the proton” [1].

The Large Hadron Collider and The Large Hadron Collider Beauty Detector

The European Organization for Nuclear Research (CERN) is a leading institution in high-energy physics research. Its excellence is primarily associated with its substantial infrastructure, including the Large Hadron Collider (LHC), the world’s largest particle accelerator, which consists of a tunnel measuring 27 kilometers and several sophisticated detectors [2]. The LHC is able to accelerate hadrons, composite particles made of quarks held together by a strong nuclear force, up to 99.9999991% of the speed of light and register data from their collisions [3].

To date, this technology has enabled the successful identification of numerous subatomic particles, with Ξ_{cc}^+ being the 80th of them [4]. The detector involved in this particular finding was the Large Hadron

Collider Beauty (LHCb). Rather than being fully enclosed around the collision site, the LHCb consists of several subdetectors distributed along a single direction. This design makes it particularly well-suited for tracking the trajectories of particles thrown forward during a collision [5].



What makes this discovery conspicuous is that Ξ_{cc}^+ was the first particle found by the LHCb detector since its renovation. The most significant change was the removal of the hardware-level trigger in favor of a fully software-based trigger system, enabling the real-time reconstruction of events at the full 40 MHz readout rate [6]. Moreover, new sensitive detector elements were installed, such as silicon-pixel (VELO) detector and Ring-imaging Cherenkov (RICH1 and RICH2) detectors [6]. Thus, the modernized LHCb demonstrates the improved discrimination of charged hadrons and enhanced capability in the identification of charm-containing hadrons. Furthermore, it now operates at a luminosity approximately five times greater than that of its original configuration.

The Nature and Detection of Xi-cc-plus

Ξ_{cc}^+ is a hadron consisting of three quarks, known as a baryon, which has a quark composition ccd, meaning that it contains two charm quarks and one down quark [1]. This specific quark arrangement makes Ξ_{cc}^+ resemblant in certain aspects of another, widely recognized baryon, the proton. The proton consists of one down quark and two up quarks. Since the down quark carries a charge of $-\frac{1}{3}$, while the up and charm quarks each carry a charge of $+\frac{2}{3}$, both the proton and the Ξ_{cc}^+ baryon have a total electric charge of $+1$.

However, there are also fundamental differences between these two baryons. Taking into consideration that the mass of a charm quark is approximately $1280 \text{ MeV}/c^2$, while that of an up quark is about $2.2 \text{ MeV}/c^2$, the Ξ_{cc}^+ is nearly four times as massive as the proton ($\approx 3619.97 \text{ MeV}/c^2$ and $\approx 938.27 \text{ MeV}/c^2$, respectively) [6].

The discovery of Ξ_{cc}^+ stems from data collected on proton-proton collisions executed in the LHC in 2024 during the Run 3 [6]. More specifically, the particle was detected through the reconstruction of its characteristic weak decay channels present in selected events registered by the LHCb detector. At the quark level, a weak decay process involves the transformation of one quark flavor into another, mediated by the weak interaction [7]. This mechanism enables transitions that are forbidden under the strong interaction, thereby allowing unstable hadrons to decay into less massive particles.

As a short-lived baryon, the Ξ_{cc}^+ does not propagate over macroscopic distances; instead, it undergoes weak decay into a set of lighter, long-lived hadrons. Accordingly, the identification process relied directly on the observation of specific final-state particles, notably the single charmed λ_c^+ baryon accompanied by K^- and π^+ mesons. These particles form the decay chain represented by the following equation [6].

$$\Xi_{cc}^+ \rightarrow \lambda_c^+ + K^- + \pi^+.$$

Invariant mass reconstruction is a method used in particle detectors to determine the

mass of an unstable particle by measuring the energy and momentum of its decay products. Using relativistic energy-momentum conservation, the invariant mass of the parent particle is calculated from the summed four-momenta of the final-state particles [8]. A peak in the resulting invariant mass distribution indicates the presence of the corresponding particle resonance. Using this method, an explicit peak was found in about 915 events—near the value of $3620 \text{ MeV}/c^2$. Local significance around the obtained signal was evaluated with a likelihood ratio test and calculated to exceed the 7σ . This indicates a level of certainty far beyond the statistical requirement, given that the threshold for declaring a discovery in particle physics is 5σ [4].

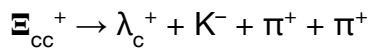
Similar Particles

Ξ_{cc}^+ is not the first doubly charmed baryon found by CERN; its precursor, Ξ_{cc}^{++} was discovered in 2017 [1]. This holds great relevance for the discovery of Ξ_{cc}^+ , as Ξ_{cc}^{++} is its isospin partner. Specifically, together the two baryons form an isospin doublet [6]. The isospin doublet refers to a pair of particles that are nearly identical in terms of strong-force interactions and masses, differing only by the replacement of an up quark with a down quark and, consequently, by their electric charge. This is exhibited by proton and the neutron, a principal example of isospin partners.

Ξ_{cc}^{++} , with a quark composition ccu, holds the elementary charge of $+2$ and has an approximate mass of $3621.4 \text{ MeV}/c^2$ [6]. This hadron held a pivotal role in constraining the expected parameters of sought Ξ_{cc}^+ . Theoretical models predicted the mass difference between these baryons to lie within the range of -0.4 to $-2.3 \text{ MeV}/c^2$ [6]. Estimates were in close agreement with experimental results; this high consistency provided strong support for confirmation of Ξ_{cc}^+ baryon's existence.

Structural differences between Ξ_{cc}^+ and Ξ_{cc}^{++} account also for more disparities in their displayed properties. In particular, Ξ_{cc}^{++} has a prominently longer lifetime, measured

by the LHCb experiment at 256 fs. In the analyses, a range of 15–160 fs is assumed for Ξ_{cc}^{++} [6]. When those lifetimes terminate, both hadrons undergo weak decays resulting in various products. The primary difference between these processes is the presence of an additional pion (π^+) in the decay products of the Ξ_{cc}^{++} baryon.



Direction of Further Research

The discovery of Ξ_{cc}^{+} baryon and knowledge of its properties hold substantial significance, offering valuable insights into ongoing and future developments in particle physics. According to LHCb spokesperson Vincenzo Vagnoni, the discovery is expected to make a major contribution to research within quantum chromodynamics (QCD) [4]. QCD is the component of the Standard Model that describes the dynamics of interactions occurring between particles carrying color charge.

Color charge is a charge analogous to electrical charge, with the exception that it is not associated with electromagnetic forces, but with strong forces instead [9]. It is carried by quarks, conventionally labelled as red, green, or blue, and by gluons—bosons which are force carriers and mediate interactions between them. The discovery of Ξ_{cc}^{+} provides a unique laboratory for testing the laws of QCD in a “heavy-quark dominated” environment. In the proton (uud composition) the constituent quark masses are sufficiently small that the majority of the proton’s mass arises not from the quarks themselves, but from their kinetic energy and the dynamic contributions of the gluon fields [1]. In

contrast, the mass of Ξ_{cc}^{+} is more directly influenced by the intrinsic masses of its constituents. This shift enhances its utility as a system for probing the mechanisms by which QCD accounts for hadron mass generation.

Furthermore, the discovery prompted scientists at CERN to search for the Ω_{cc}^{+} particle—a hadron that is also a doubly charmed baryon, but forms an isospin singlet [6]. It contains a strange quark, giving it the ccs quark composition. The detection of this baryon was announced at the Beauty 2026 conference in Maastricht [10]. Analogously to the discovery of Ξ_{cc}^{+} , the observation of Ω_{cc}^{+} is based on Run 3 data collected with the upgraded LHCb detector in 2024. Scientists at CERN also adopt a long-term perspective, suggesting that after the second upgrade of the LHCb detector, the observation of a triply charmed baryon may become experimentally accessible [6].

Conclusion

Overall, the observation and study of the Ξ_{cc}^{+} baryon constitute a significant milestone in physics. The expanding body of knowledge surrounding doubly charmed baryons opens new avenues for both experimental and theoretical inquiry, including the exploration of multi quark systems, improved understanding of strong interaction dynamics, and the potential discovery of previously unobserved states. In light of the planned improvements of the LHCb detector, further groundbreaking discoveries in this field remain highly anticipated.

By Maja Madej

References:

- [1] B. Pietrzyk, “Observation of the doubly charmed heavy proton Ξ_{cc}^{+} ,” LHCb Outreach, March 17th, 2026.
- [2] “The Large Hadron Collider,” CERN. Retrieved 20 Apr 2026.
- [3] A. Dauvergne, “The LHC in numbers,” CERN Document Server.
- [4] “LHCb Collaboration discovers new proton-like particle,” home.cern.
- [5] “LHCb,” CERN.
- [6] S. Han, “Search for the doubly charmed baryon Ξ_{cc}^{+} with the upgraded LHCb detector,” March 16th, 2026.
- [7] T. Mannel, “Weak decays of heavy quark systems,” Theory Division, CERN, 1994.
- [8] M. Barnett, E. Etzion, “A Search for New Physics Processes using Dijet Events,” ATLAS Experiment. June 21th, 2011.
- [9] O. W. Greenberg et al., “Color Charge Degree of Freedom in Particle Physics,” Compendium of Quantum Physics, pp. 109–111
- [10] B. Pietrzyk, “Observation of the doubly charmed baryon Ω_{cc}^{+} ,” LHCb Outreach, June 3rd, 2026.

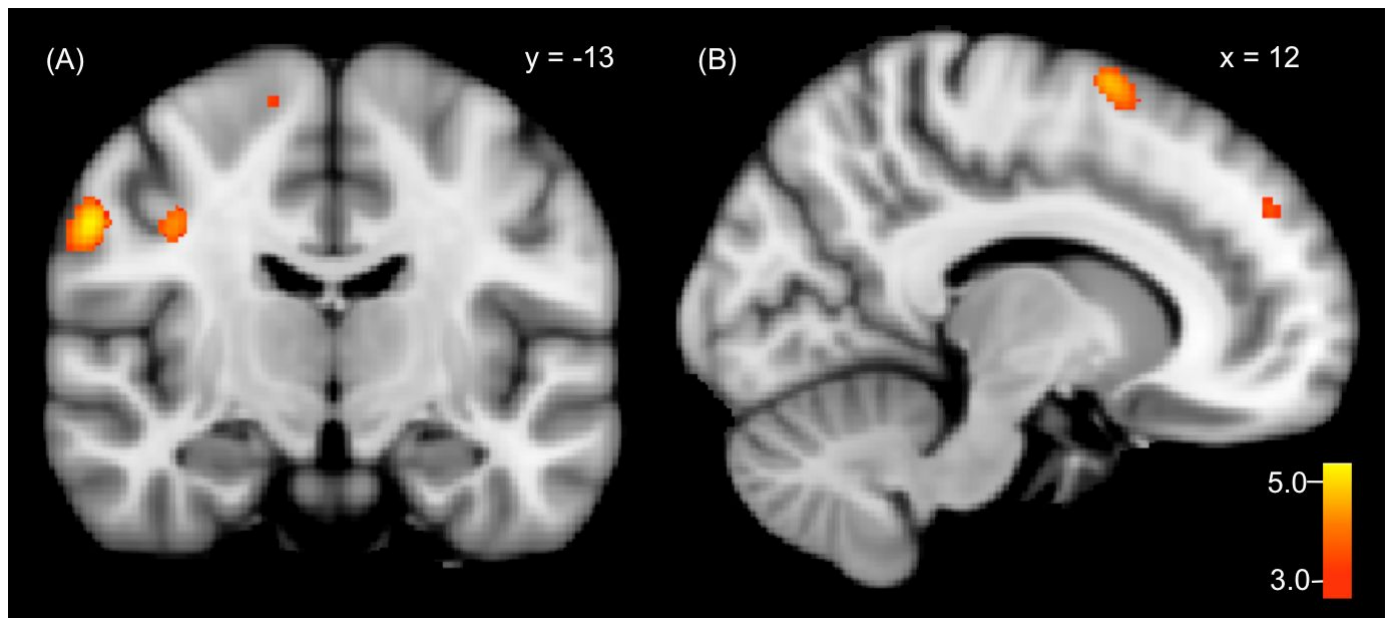


Figure 1: fMRI scan of DID patient [1].

Investigating the Fractured Self: The Neuropsychology of Dissociative Identity Disorder.

Introduction

A sense of self is often assumed to be stable, continuous and personal. However, in certain psychological circumstances, this idea begins to unravel. Reported estimates suggest that approximately 1-5% of the international population are affected by a dissociative disorder, including Dissociative Identity Disorder (DID) and Depersonalisation-Derealisation Disorder (DDD) [2]. Dissociation itself, is characterised by an involuntary disruption in an individual's thoughts, identity, memory and awareness, affecting the individual's perception of self and consciousness [3]. While dissociation is often pathologised using a clinical approach, there are other profound aspects to it which may reside in a person's reshaping of the psyche through psychological responses to prolonged trauma.

How It Begins

An onset of Dissociative Identity Disorder occurs at a young age, ranging from 5-9 years old, while the brain is still in its late pre-operational early concrete operational stage and a cohesive sense of self is still maturing. The disorder itself can develop as a result of chronic trauma in childhood,

which causes the brain to compartmentalise these traumatic experiences to isolate the stress. Ultimately, this serves as an adaptive response to protect the individual from further suffering and exposure to psychological distress [4]. Rather than the mind literally splitting into an abundance of alternate selves during dissociation, DID can be better understood as the brain's failure to consolidate one unified identity, which causes an eventual emergence of multiple identities as distinct, personal states of mind to manifest over time [5].

The Lived Experience

Symptoms can generally differ from person to person; however, the most common are memory gaps, known as dissociative amnesia for some. This is seen to have a correlation with a smaller and typically impaired hippocampus [1]. As the hippocampus plays a key role in the consolidation of memory, this impairment may explain memory fragmentation, and difficulties in memory integration [1]. Additionally, individuals may experience having an unstable perception of self, which is thought to be caused by a dysregulated Default Mode Network, which is responsible for self-referential thought and introspection. This is then typically accompanied by a hyperconnectivity between the DMN and the Frontoparietal Network, which handles executive functioning [6].

Since the DMN is responsible for self-referential processing while the Frontoparietal Network helps direct attention towards tasks, excessive communication between the two may blur the distinction between internally generated experiences and the outside world [6]. This can be associated with a struggle to properly integrate internal thoughts with external reality and “trance-like” states where individuals may seem unresponsive to external stimuli [6]. Furthermore, some patients may find themselves in a state of “passive influence”, where headmates who are not currently prominent tend to intrude on the altar in front.

Interestingly, some patients may report a difference in abilities between identity states which usually would not be possible for the alter alone [7]. The alter in front could also experience residual influence, where their thoughts and abilities are unconsciously warped by the recently departed alter – this ranges in potency depending on the former alter.

The Neuroscience of a “Switch”

In DID, the process of switching refers to the transition between identity states, where one becomes more prominent while the other becomes less active. They can be forced, triggered by stress, trauma or specific sensory inputs or mutually agreed [8]. Common indicators of an individual switching may be depersonalisation, disorientation or episodes of fugue. Externally, some cues may include changes in demeanor, visible perplexity and potentially even muscle spasms [9]. These switches can range between being gradual and overt, to being incredibly seamless and discrete.

Functional neuroimaging (fMRI) studies reveal a decreased activity in the prefrontal cortex, which is responsible for control and self-awareness, and an altered activity in the amygdala [9]. Moreover, alters may exhibit distinct cognitive states, which become more prominent during switching, allowing access to different abilities or

memories that may not otherwise be available when not in front.

Common Misconceptions

Prior to DID becoming a formally recognised disorder within the ICD, it was often mistaken for hysteria [10]. Nowadays, it may be misdiagnosed as Schizophrenia and Borderline Personality Disorder due to an overlap in symptoms [11]. Nevertheless, contemporary psychology still demonstrates persistent misconceptions surrounding the disorder. Brand et al. outlines and debunks six empirical outdated myths about DID, ranging from the legitimacy of DID and its causes to the effectiveness of treatment on the patient [12].

Despite DID being classified within countless diagnostic manuals, remnants of doubt regarding its authenticity still loom within society. Claims that it is iatrogenic and caused as a result of a self-fulfilling prophecy, typically deviate from the since disproved fantasy model of dissociation [13]. This model proposed that DID develops from suggestibility and sociocultural influences, rather than as a response to excessive psychological trauma, this idea has since been disproved and displaced by the trauma model of dissociation [13].

Within the media, it is not uncommon to see DID representation follow a “Jekyll-and-Hyde” like trope, with caricatural “malicious alter-egos” taking control while the individual remains in a passive, dreamlike repose. Depictions like these have served as a significant contributor to public stigma as they obscure the complex nature of the disorder in favour of dramatic narratives, ultimately reinforcing negative stereotypes about DID patients as hysterical and erratic [14].

Conclusion

To conclude, Dissociative Identity Disorder reflects the mind’s capacity for self-preservation and adaptability in times of crisis. It reflects a profound and core idea

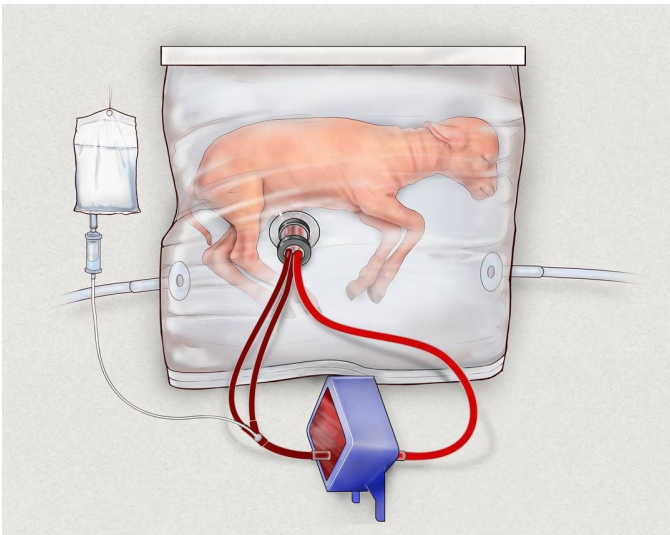
that identity is not fixed, but may become fragmented as a means of self-preservation at the cost of self-continuity. Despite often being sensationalised in the media through stereotypical and exaggerated portrayals, DID offers a profound insight into the

plasticity of the human mind. Ultimately, the intricate relationship between psychological experience and cognitive function gives a more holistic understanding of the human psyche.

By Maria Wahba

References:

- [1] Y. R. Schlumpf et al., *"Dissociative Part-Dependent Resting-State Activity in Dissociative Identity Disorder: A Controlled fMRI Perfusion Study,"* 2014.
- [2] P. Mitra and A. Jain, *"Dissociative Identity Disorder,"* in StatPearls. Treasure Island, FL, USA: StatPearls Publishing, 2023.
- [3] E. Cardeña, "The Domain of Dissociation," in *Dissociation: Clinical and Theoretical Perspectives*, pp. 15–31, 1994.
- [4] V. Şar, M. J. Dorahy, C. and Krüger, "Revisiting the etiological aspects of dissociative identity disorder: a biopsychosocial perspective," *Psychology research and behavior management*, pp.137-146, 2017.
- [5] D. Rothschild, *"On Becoming One-Self: Reflections on the Concept of Integration as Seen Through a Case of Dissociative Identity Disorder,"* *Psychoanalytic Dialogues*, pp. 175–187, 2009.
- [6] V. Menon, "Dissociation by network integration," *American Journal of Psychiatry*, 178(2), pp.110-112, 2021.
- [7] F. W. Putnam, "VII Dissociation and Disturbances of Self," *Disorders and dysfunctions of the self*, 5, p.251, 1994.
- [8] G. E. Tsai et al., "Functional magnetic resonance imaging of personality switches in a woman with dissociative identity disorder," *Harvard review of psychiatry*, 7(2), pp.119-122, 1999.
- [9] R. L. Savoy et al., "Voluntary switching between identities in dissociative identity disorder: A functional MRI case study," *Cognitive Neuroscience*, 3(2), pp.112-119, 2012.
- [10] B. Foote, "Dissociative identity disorder and pseudo-hysteria," *American journal of psychotherapy*, 53(3), pp.320-343, 1999.
- [11] N. A. Ashikin, "Unmasking the Mind: A Journey Through Misdiagnosis to the True Identity of Dissociative Identity Disorder," *BJPsych Open*, 11(S1), pp.S292-S293, 2025.
- [12] B. L. Brand et al., "Separating fact from fiction: An empirical examination of six myths about dissociative identity disorder," *Harvard review of psychiatry*, 24(4), pp.257-270, 2016.
- [13] C. J. Dalenberg et al., "Evaluation of the evidence for the trauma and fantasy models of dissociation," *Psychological bulletin*, 138(3), p.550, 2012.
- [14] B. L. Snyder et al., "It's not just a movie: perceived impact of misportrayals of dissociative identity disorder in the media on self and treatment," *European Journal of Trauma & Dissociation*, 8(3), p.100429, 2024.



Clot to Trot: Longevity of Localised NOSA in Artificial Placentas Across a Variety Of Genotypes in Lamb Models.

Introduction

Artificial placentas (APs) offer a promising intervention for extremely preterm infants, but their thrombogenicity necessitates anticoagulation therapy, which in turn increases bleeding diathesis in an already vulnerable population [1]. Nitric oxide surface coating (NOSA) presents a localised alternative, using a slow-release polymer layer on the AP interior to disperse nitric oxide directly onto passing blood, suppressing platelet activation and aggregation without systemic anticoagulant exposure [2], [3]. However, the long-term safety and efficacy of NOSA remain poorly understood [2].

This study addresses that gap using an ovine model, selected for its cardiovascular similarity to the human system [3]. Healthy, γ -glutamyl carboxylase deficient, and haemophilia A lambs will be divided into six groups of ten by genotype and treatment type (heparin versus NOSA), and monitored over three months across thrombogenicity, bleeding diathesis, and general health outcomes. The aim is to establish the long-term safety profile of NOSA and advance APs toward clinical use in the treatment of preterm infants.

Related Scientific Background

Artificial placenta (AP) systems have the

potential to save the lives of infants by significantly increasing their odds of survival [1]. While not currently used in clinical settings, they would be used on individuals that are failing maximal medical therapy and save them by replicating the typical functions of the mother's body and placenta [1]. However, the artificial material may trigger major cellular activation of platelets upon contact with blood, causing extreme thrombogenicity [2], [3]. Systemic anticoagulants, such as Heparin, limit the incidence of thrombosis, but result in an increase in bleeding risk that can lead to fatal hemorrhage throughout the body [3]. Premature newborns have a higher baseline rate of intraventricular hemorrhage that complicates balancing the risk between clotting and bleeding in treatment [3].

As a result, NOSA is being studied as a non-systemic anticoagulant that would only affect a localised region due to its short half-life intravascularly [4]. Studies have shown the effectiveness of NOSA in the short term of up to 7 days, but its long-term effects on individuals of varying genetic health has not been adequately researched, leading to uncertainties such as coating durability, systemic effects, and long-term effectiveness [3].

Testing the effects of NOSA on certain genetic disorders would shed light on how it affects individuals with increased bleeding risk, ensuring that the treatment remains safe for all in application. Hemophilia A—a congenital disorder that leads to increased bleeding and bruising—in sheep demonstrates effects such as soft tissue hemorrhage and hemarthrosis [5]. Γ -glutamyl carboxylase deficiency in sheep mirrors the human condition of vitamin K-dependent coagulation factors (VKCFD), as both result in prolonged prothrombin time and activated partial thromboplastin time [6].

Research Idea

The proposed research idea is a study on the effects of NOSA use in APs supporting premature lambs of varying genetic health,

in relation to bleeding diathesis, over several months. Lambs are selected as the model organism due to their well-documented cardiovascular and hematological similarity to humans, making physiological responses observed in lambs reasonably transferable to human neonates [3]. Additionally, the availability of lambs with naturally occurring bleeding disorders such as hemophilia A and γ -glutamyl carboxylase deficiency makes them uniquely suited to studying how NOSA performs across genetically variable subjects [6].

Rationale

This proposal studies how NOSA performs over extended periods of time to reflect the conditions that NOSA would be used in for a clinical setting. Furthermore, it clarifies how it may affect infants with increased susceptibility towards bleeding diathesis differently, compared with healthy infants. Another aspect of study will be the durability of the NOSA itself, as there is the possibility it may lose efficacy over longer periods of time. With more information with regard to these subjects, there will be more knowledge on the long-term feasibility of NOSA. Once in use, individuals will likely rely on the support of an AP for several months at a time, making this crucial to understand [7].

Methodology and Research Approach

The subjects consist of three types of lambs: healthy, γ -glutamyl carboxylase deficient, and hemophilia A. Groups will include one control group that is treated with unfractionated heparin (UFH), and one NOSA treated group for each type of lamb (n=60), totalling to 6 groups of 10 testing individuals each. Groups of 10 were chosen to allow for careful monitoring of each subject, while reducing the potential for variability to confound results and minimizing the impact of unexpected adverse events on overall study outcomes. This expands on the Major et al. [8] paper by increasing group size and variability so the results will apply better to a larger

population.

All lambs will be delivered prematurely via cesarean section at 85 days. AP circuits will use jugular vein drainage and umbilical vein reinfusion to mimic natural fetal blood flow and place low strain on the heart, similar to the setup by Fallon BP et al. [1].

NOSA-treated groups will receive NOSA coating on all circuit components that come into contact with blood. The NOSA coating will contain a hydrophobic polyurethane base underneath a layer made out of nitric oxide donors (DBHD/N₂O₂), and a biodegradable polymer (PLGA) that will allow for sustained release of the anticoagulant. This experiment will run continuously for three months, where recordings of the lamb's health will be taken every hour. Furthermore, visual observations will be taken of the AP circuit to monitor the potential formation of thromboses. A time period of three months was chosen to simulate the realistic clinical duration infants may require support from an AP, as well as giving time for delayed complications and long-term trends to develop.

Three main statistics will be monitored: circuit thrombogenicity, bleeding status, and general health. Increases in circuit pressures, requiring of a component change, or documented clot events will alert to increased circuit thrombogenicity. These three indicators can be weighted and combined into a composite thrombogenicity score, where minor pressure elevations contribute least, component changes contribute moderately, and confirmed clot events contribute most severely. A standardised scoring will be performed to compare performance across groups. Daily assessments will be performed on each subject to check for symptoms of excessive bleeding, such as external bleeding at cannulation sites, mucosal bleeding, and soft tissue hematomas now. There is no universally-accepted scoring system for bleeding in lambs, so this experiment will rely on descriptive reporting and parameters.

General health records will be kept of days

survived and bodyweight. Heart rate and oxygen saturation will be checked hourly. These factors help ensure long-term feasibility of the AP and NOSA, by indicating whether the subject is responding safely to the treatment. Blood samples will be taken weekly to check for platelet count and fibrin deposition. The weekly timeframe minimises stress on the subjects while still allowing for the observation of any unexpected trends that develop slowly. Data analysis will be performed using a two-way ANOVA, with treatment type and genotype as factors, to demonstrate the difference in behaviour and outcome of each group.

Conclusion

This paper proposes a three-month longitudinal study examining the longevity and anticoagulant effectiveness of NOSA-coated AP circuits across healthy, hemophilia A, and γ -glutamyl carboxylase

deficient lamb models, using thrombogenicity, bleeding diathesis, and general health as primary outcome measures. This study aims to clarify whether NOSA remains a safe and durable alternative to systemic anticoagulation over clinically realistic timeframes. However, lamb models cannot perfectly replicate the complexity of an extremely premature human infant, and a sample size of ten per group may limit the generalisability of findings.

The absence of a universally accepted bleeding scoring system for lambs also introduces some subjectivity into one of the study's core metrics. Despite these limitations, if NOSA demonstrates sustained safety and effectiveness across genetically diverse subjects, neonatologists will have a meaningfully safer tool at their disposal. NOSA treatments remove some of the risk between clotting and bleeding that currently defines care for premature infants.

By Marshall Zhang

References:

- [1] B. P. Fallon and G. B. Mychaliska, "Development of an artificial placenta for support of premature infants: narrative review of the history, recent milestones, and future innovation," *Translational Pediatrics*, vol. 10, no. 5, pp. 1470–1485, 2021.
- [2] D. G. Blauvelt, E. N. Abada, P. Oishi, and S. Roy, "Advances in extracorporeal membrane oxygenator design for artificial placenta technology," *Artificial Organs*, vol. 45, no. 3, pp. 205–221, 2021.
- [3] B. P. Fallon et al., "Extracorporeal life support without systemic anticoagulation: a nitric oxide-based non-thrombogenic circuit for the artificial placenta in an ovine model," *Pediatric Research*, vol. 95, no. 1, pp. 93–101, 2024.
- [4] D. D. Thomas, X. Liu, S. P. Kantrow, and J. R. Lancaster, "The biological lifetime of nitric oxide: implications for the perivascular dynamics of NO and O₂," *Proceedings of the National Academy of Sciences*, vol. 98, no. 1, pp. 355–360, 2001.
- [5] J. N. Lozier and T. C. Nichols, "Animal models of hemophilia and related bleeding disorders," *Seminars in Hematology*, vol. 50, no. 2, pp. 175–184, 2013.
- [6] J. S. Johnson, B. A. Soute, C. S. Olver, and D. C. Baker, "Defective γ -glutamyl carboxylase activity and bleeding in rambouillet sheep," *Veterinary Pathology*, vol. 43, no. 5, pp. 726–732, 2006.
- [7] H. Usuda et al., "Artificial placenta technology: history, potential and perception," *Placenta*, vol. 141, pp. 10–17, 2023.
- [8] T. C. Major et al., "The effect of a polyurethane coating incorporating both a thrombin inhibitor and nitric oxide on hemocompatibility in extracorporeal circulation," *Biomaterials*, vol. 35, no. 26, pp. 7271–7285, 2014.

Dirac Sea and Quantum Brain Dynamics: A Science Based Insight of Psychic Abilities

Introduction

While traditional psychology focuses on observable behavior, the underlying mechanisms of perception may be rooted in the quantum world. Exploring how negative energy particles might shape the way we experience reality bridges the gap between quantum physics and neuroscience. By hypothesizing a biological system where quantum mechanics and neural activity intertwine, quantum brain dynamics seeks to redefine our understanding of the mind-matter connection.

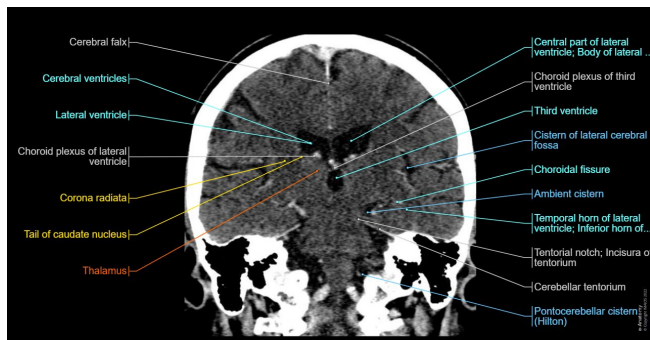


Fig 1: Anatomy of the brain [1].

Dirac Sea

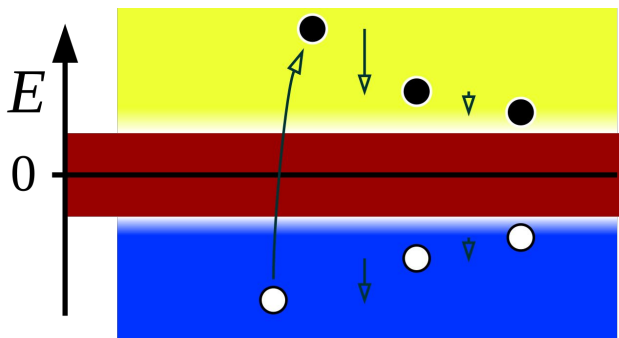


Figure 2: Neuroscientific implications.

Postulated by the British physicist Paul Dirac in 1930, the Dirac sea is a theoretical model of the electron vacuum as an infinite sea of electrons with negative energy, now identified as positrons, which are the antimatter counterparts to electrons [2]. This model suggests that the perceived vacuum between electrons is actually a dense field of positrons. However, these positrons generally exist outside observation or control of humans within classical frameworks. Hence, the vacuum is not truly empty, rather it is a dense field of

positrons.

This theory was postulated to explain the anomalous negative-energy quantum states predicted by the relativistically correct Dirac equation for electrons. The positron was originally conceived of as a hole in the Dirac sea, before its experimental discovery in 1932.

Quantum Brain Dynamics

Quantum Brain Dynamics (QBD) is a theoretical framework suggesting that the brain's cognitive and conscious functions are established by quantum mechanical processes. This framework challenges classical neuronal activity models. It suggests that phenomena such as memory and awareness are not merely the result of electro-chemical impulses, but rather involve deeper quantum interactions. Quantum Brain Dynamics offers a novel perspective on how the brain processes information, particularly in relation to higher-order functions like consciousness. Classical models of cognition and neural processing rely primarily on electrical impulses and chemical signalling. While these are well-established, they struggle to explain certain aspects of conscious experience, including the subjective quality of awareness and the unity of perception [3].

Furthermore, QBD suggests that the brain's memory storage is not limited to synaptic connections, but also involves the "ordering" of quantum states, a process where the brain maintains specific energy alignments to store information. Giuseppe Vitiello argues that memory results from the coherent condensation of quanta within the brain's water matrix: a process where water molecules, particularly within cellular structures like microtubules, spontaneously shift from a disordered (incoherent) state into a highly ordered, synchronized and macroscopic quantum state [4]. This coherent condensation involves dipole wave quanta, which acts as the medium through which the brain stores and retrieves information. In this context, the Dirac Sea serves as a reservoir of potential energy

states from which conscious “excitations” emerge [4].

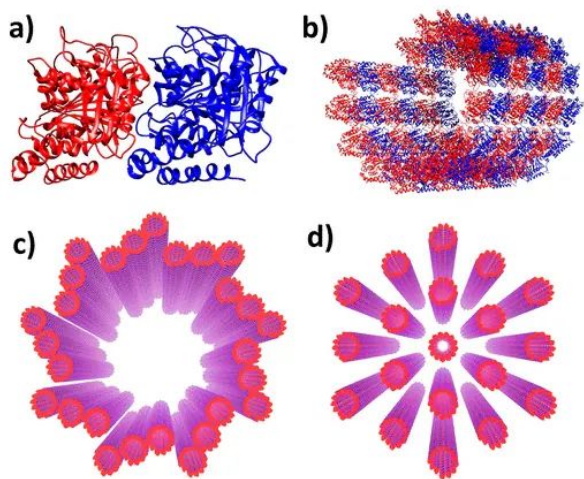


Figure 3: Protein Architectures [5].

While the exact mechanism by which anaesthetics cause a loss of consciousness remains unknown, evidence suggests they interact with microtubules [5]. The evidence of new research suggests a potential link between microtubule behaviour and quantum physics during anaesthesia.

The “Science” Behind Psychic Abilities

In a hypothetical structured system of the world where the framework of QBD is proven to exist and the Dirac Sea model is actively interacting with our reality, Quantum Entanglement and non-locality can be used as the factual foundation behind the functional reality of consciousness.

In the sense of non-local realism, If the brain operates as a quantum system, it could theoretically allow for entanglement between two individuals. In quantum physics, entangled particles share states

instantly across any distance. This idea suggests the potential for allowing “mind control” ability, merging two consciousness under one identity. Additionally, the Dirac Sea is hypothesized as a universal medium, also known as a “Zero Point Field”, through which psychic abilities would be the act of “sampling” the field’s information density. However, the primary “scientific” hurdle for telepathy is the no-communication theorem, which states that entanglement cannot be used to transmit information faster than light.

Studies conducted by the Global Consciousness Project at Princeton suggest Random Event Generator (REG) devices that produce unpredictable sequences, have been used to study the influence of human intention on physical systems. The aim was to see if human intention affects physical machines. While controversial, their data suggests that during major world events, the “entropy” of these machines decreases, implying a non-local field of human influence. The connection with reality lies in the Quantum Field Theory (QFT) approach to consciousness.

Conclusion

Psychic abilities can be potentially accessed through the Dirac Sea to manipulate the “endless possibilities” of quantum wave functions. A detailed exploration of wave functions is omitted here to maintain focus on the neuroscientific application. This process can be viewed as an equation of energy transfer involving waves.

By Omar Faruque

References:

- [1] A Micheau, D Hoa, “Normal cranial CT Scan of the head: brain, bones of cranium, sinuses of the face,” IMAIOS, Jul. 12, 2015.
- [2] “Dirac Sea,” Wikipedia, 2026.
- [3] J. Benardo, A. Kemi “Quantum Brain Dynamics,” ResearchGate, 2025.
- [4] G. Vitiello, *My Double Unveiled: The Quantum Physics of Consciousness*. Amsterdam, Netherlands: John Benjamins Publishing Company, 2001.
- [5] N. S. Babcock G. Montes-Cabrera K. E. Oberhofer M. Chergui G. L. Celardo P. Kurian “Ultraviolet Superradiance from Mega-Networks of Tryptophan in Biological Architectures”, *J. Phys. Chem. B* 2024, 128, 17, 4035–4046

inflammation, improved wound healing, and overall skin rejuvenation [5]. Effects on aging are often most noticeable through its effects on skin and hair. Some evidence suggests that red light therapy may help improve both by widening blood vessels, which allows more blood and nutrients to reach skin cells and hair follicles. This can support better cell function and repair over time, leading to improvements in skin texture and hair quality. While these changes may create a more youthful appearance, they do not necessarily mean that aging is “being slowed”—just that some of its visible effects are being reduced [6].

However, when considering the effectiveness of red light therapy, several other factors should be considered. Genetics play a key role in determining outcomes, influencing both the timeline and overall effectiveness of results. As a result, genetics may affect how well your cells respond to the treatment [5], [6].

Additionally, it is reasonable to expect that individuals who spend as much money on devices like these are also more likely to invest more time and effort into other health-related habits, such as diet, skincare routines, or other wellness practices. In other words, those who spend a lot of money on a device intended for self improvement are also more likely to have more resources to use other evidence-backed routines to help achieve a common goal.

Limitations and Misconceptions

One of the biggest limitations of red light therapy is the cost. Even though it offers many potential benefits, devices can be expensive, usually ranging from about \$300 to over \$1000 USD [2], [4]. Cheaper options are available, but they can come with risks. Lower-cost devices may use low-quality LEDs, produce the wrong wavelengths, or fail to deliver enough energy to be effective. In some cases, incorrect wavelengths may cause unwanted effects on the skin or reduce the effectiveness of the treatment. Since red light therapy depends on a variety of factors to yield desirable results, good quality devices play a key role [3], [4].

Furthermore, as with many other treatments, results take time. People often become overly excited at first, but the disappointment may be much more noticeable in comparison just after a week when no visible changes appear.

Conclusion

Whether red light can slow aging depends on several factors. According to current evidence, the theory of “stopping aging” is not true [7]. However, the biological processes associated with aging can be slowed and their visible effects reduced through red light therapy [1], [3], [4], [5]. By stimulating mitochondria and increasing cellular energy reproduction, red light therapy may support processes such as skin repair, hair growth, and reduced inflammation.

By Pepper Koh

References:

- [1] C. Pagán, “Red Light Therapy: What is it?,” WebMD.com.
- [2] “Light Therapy Market Size to Reach USD 1.79 Billion by 2033, Expanding at a CAGR of 5.50% - SNS Insider,” finance.yahoo.com, Dec. 18, 2025.
- [3] “Red Light Therapy,” my.ClevelandClinic.org.
- [4] T. Breillat, “Why Your Cheap Red Light Therapy Device Isn’t Working (And What to Do About It),” healthlightllc.com.
- [5] D. Moon, “Red Light and Photobiomodulation: ATP from Photons,” geneticlifehacks.com.
- [6] H. Armitage, “Red light therapy: What the science says,” med.stanford.edu.
- [7] M. Berghoff, “Can ageing be slowed down?,” age.mpg.de.



The Return of The Dire Wolf

Introduction

There are three specimens, Remus, Romulus and Khaleesi who are often referred to as dire wolves. While the howl of this extinct animal has not been heard for over 10,000 years, the biotechnology company Colossal claims that these three animals are the world's first de-extinct dire wolves. Conversely, some scientists say that the animals are categorized as genetically modified gray wolves.

The Dire Wolf

The dire wolf, *Canis (Aenocyon) dirus*, was the largest of the North American late Pleistocene canines. Its skull measured as long as 12 inches, and its teeth were bigger and stronger than those of modern gray wolves, identifying them as carnivorous mammals. Horses made up a key prey species, while sloths, bison, and camels made up a minority of their diet [1].

Most scientists, as reported by The Conversation in an article, agree that the extinction of the dire wolf about 10,000 years ago was simultaneous with the extinction of most of their prey. Dire wolf fossils were found in North and South America. In the continent of North America, they have been identified as far north as Alaska to southern Mexico. As noted in [2], numerous dire wolf skeletons have been exposed in the Americas. Overall, the dire wolf seemed to be more heavily built than agile, with greater muscle and a heavier build than other Pleistocene canids or

those that still occupy the planet today. However, their greater size, based on fossils, did not compromise their ability to travel across diverse landscapes, pursue large prey, or coexist peacefully with other North American fauna. In fact, dire wolves were effective top predators because of their physique, not in spite of it, [3].

Somatic Cell Nuclear Transfer

To understand the comeback of the dire wolf, first it is necessary to understand somatic cell nuclear transfer (SCNT), which is a cloning technique where the nucleus of an egg cell is removed to create an enucleated cell. Then, a somatic cell, which is any non-sex cell, is taken out of the animal to be cloned. The nucleus of the somatic cell is placed inside the egg which has had its own genetic code removed. The placement is either done by direct injection or fusion using an electric charge, thereby initiating embryonic development.

The egg is then caused to split, thus forming an embryo, which is then implanted in a surrogate mother's womb. If the procedure goes as expected, the embryo develops normally and the clone is born.

Even though this method has cloned several different animals like sheep, cows, mules, rabbits, and dogs, few are actually born, there is a low survival rate of merely making it to birth. Those cloned animals that do survive will age and die prematurely as their DNA is taken from adult cells that have undergone telomere shortening, the loss of protective ends of the chromosomes with each cell division, as a normal process of aging [4].

The Process Of Bringing Back The Dire Wolf

Supported by Colossal Biosciences, the Dire Wolf Project uses gene editing to restore the ancient predator. Scientists analyzed the genome of dire wolves from a fossilized tooth in Ohio and a 72,000-year-old skull in Idaho. They

compared the genomes to that of the gray wolf and identified 20 differences in 14 genes.

They cultured endothelial progenitor cells (EPCs), which constitute the lining of the blood vessels, from living gray wolves, then had 14 genes in the nuclei edited to match the genes of the dire wolf using clustered regularly interspaced short palindromic repeats (CRISPR) technology. This works with RNA to find specific DNA sequences and Cas9, an endonuclease that cuts DNA sequences. Once the genetic material is divided, scientists can delete, insert or repair genes. After editing the cells to display the characteristics of the dire wolf, they used the cloning technique somatic cell nuclear transfer: the reedited cell nuclei were inserted into egg cells whose native nuclei had been excised, developing into healthy embryos that carried the dire wolf genetic profile.

As explained in [5], genes are likely to have more than one consequence. The dire wolf has three genes that code for its light colored coat, but in gray wolves, the identical genes can lead to deafness and blindness. To avoid this, Colossal's research team hacked two other genes to disable black and red pigments, allowing for the characteristic light color of the dire wolf without causing harm to the gray wolf.

The de-extinction process requires substantial funding, though Colossal has an estimated evaluation of \$10.2 billion. In addition, the company is not alone. It is partnering with conservation organizations such as the American Wolf Foundation, the Mauritian Wildlife Foundation, Save the Elephants, and Conservation Nation. Colossal also partnered with the Indigenous MHA Nation tribes (Mandan, Hidatsa, and Arikara) to collaborate on the project, and the tribes have been interested in having dire wolves run on their territories in North Dakota [5].

The Birth of The Dire Wolf

On Oct. 1, 2024, the surrogates each gave

birth to 2 males, Romulus and Remus. On Jan. 30, 2025, a female pup, Khaleesi was born. The surrogate hosts were maintained during gestation at Colossal's animal-care center, where they were monitored frequently and had weekly ultrasounds from company scientists and vets. The three pups were all delivered by timed cesarean section to minimize complications at birth [2].

The deliveries were performed by a team of 4 surgery personnel while 4 other attendants cared for the mothers and infants as they recuperated from anesthesia. The infants were nursed by the surrogate for a short period of days, after which the Colossal team removed them from the host mother and bottle-fed them [2], [6].

Controversy and Skepticism

Some independent researchers assert that Romulus, Remus, and Khaleesi are not authentic dire wolves. University of Otago zoologist Philip Seddon in New Zealand explained the animals are "genetically modified grey wolves." Researchers indicated considerable biological differences between contemporary dire wolves and the dire wolves that lived in the Americas 10,000 years ago [7].

Paleogeneticist Dr. Nic Rawlence has explained that ancient DNA within fossilized bones is too degraded to biologically reproduce or clone. Dr. Rawlence argues that adjusting only 14 genes from 19,000 does not create a true dire wolf, but a gray wolf with characteristics resembling those of the extinct animal.

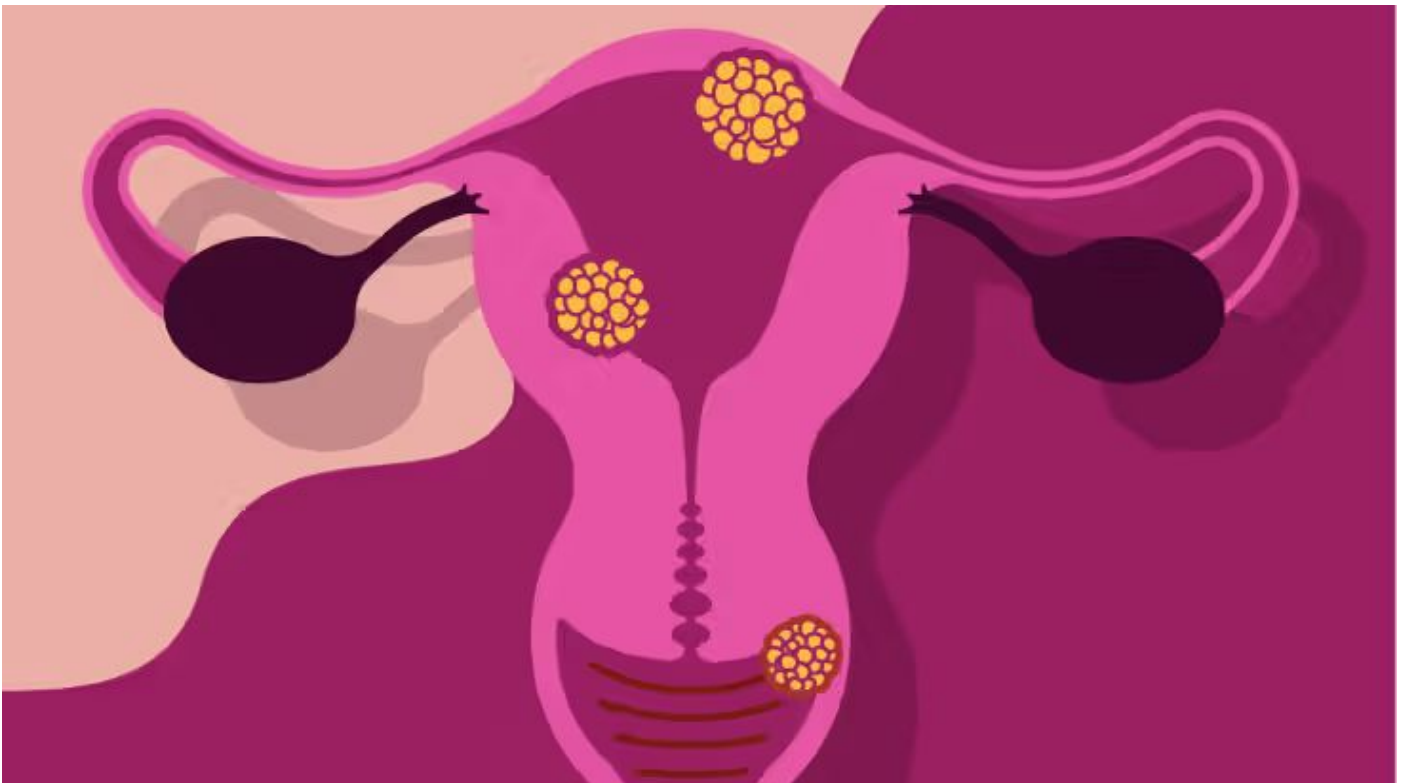
Conclusion

Despite the debate among scientists, the three wolves represent a breakthrough in gene technology and the de-extinction process through the integration of advanced biotechnology with the DNA of extinct species.

By Rim Laasri

References:

- [1] National Park Service, "[Dire wolf](#)," U.S. National Park Service.
- [2] J. Kluger, "[The return of the dire wolf](#)," TIME, Apr. 2025.
- [3] Colossal Biosciences, "[Dire wolf](#)," Colossal Biosciences.
- [4] JoVE Science Education, "[Reproductive cloning – Somatic cell nuclear transfer](#)," JoVE.
- [5] J. Kluger, "[The science behind the return of the dire wolf](#)," TIME, Apr. 2025.
- [6] J. deJong, "[What is a dire wolf? How formerly extinct species compares to gray wolves](#)," Newsweek, Apr. 2025.
- [7] V. Gill, "[Experts dispute claim dire wolf brought back from extinction](#)," BBC, Apr. 2025.



Bacterial Vaginosis Treatments and How They Can Change Vaginal Microbiota Composition

Introduction

A healthy vaginal microbiota is typically colonized by *Lactobacillus* which maintains bacterial diversity low and prevents infections [1]. Changes to this fragile balance may cause dysbiosis and infections such as Bacterial Vaginosis (BV). Although BV is treatable with antibiotics, recurrence is common for reasons not completely understood. Some risk factors are: lack of condom use, more than one sexual partner, smoking, vaginal douching, and incomplete antibiotic treatments. There are many gaps in research in women's health, particularly in the vaginal microbiota and the mechanisms by which it maintains itself, differs between races, and is inherited or developed. Understanding the way bacterial vaginosis treatment can change vaginal microbial composition is essential to determine its implications for recurrence.

Bacterial Vaginosis: Symptoms and Diagnosis

The vaginal microbiota is characterized as the least diverse niche, or area, in the human body. Where *Lactobacillus*

colonization is widely regarded as the most protective composition in a healthy vagina [1]. Shifts to these bacterial communities may influence the health of the vaginal environment by increasing susceptibility to infection and risk of cervical disease [2]. Specifically, changes may result in an infection like BV, which is characterized by a depletion of *Lactobacilli* and an overgrowth of anaerobes such as *Gardnerella vaginalis*, *Ureaplasma urealyticum*, *Prevotella* sp., and *Atopobium vaginae* [3].

The way BV is diagnosed can be varied, and is highly dependent on time and resources available to the patient. On one hand, the Amsel criteria may be used to diagnose BV clinically if the patient presents with at least three symptoms from a list of criteria: a vaginal pH greater than 4.5; thin, gray or yellowish vaginal discharge; amine odor when adding a 10% potassium hydroxide solution to the sample; vaginal cells covered in bacteria, known as clue cells, in more than 20% of the total cell population [4]. Whereas, the Nugent Score may be used to visualize and quantify the stained bacterial sample under oil immersion microscopy [4], [5].

Oil immersion microscopy functions by using immersion oil in the space between

the lens and the slide, which directs light in a straight line, resulting in high magnification and clear visualization of the samples. The Nugent Score is considered the gold standard, but it is a more lengthy process that requires specialized microscopy skills and the appropriate equipment in the clinical setting. As such, the Amsel criteria is more widely used.

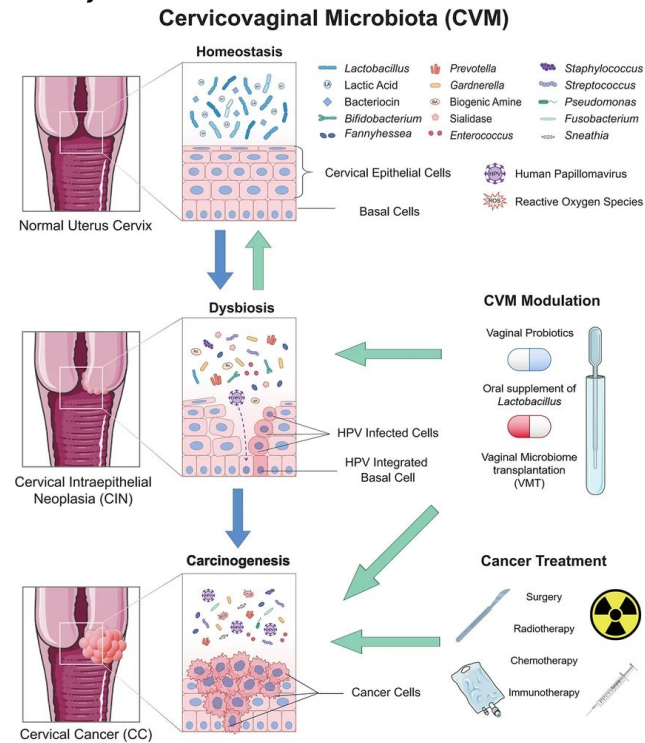


Figure 1: Typical cervicovaginal microbiota and the changes the epithelial cells may undergo once this composition shifts and become less protective [2].

Treatment and Clinical Study

Treatment for BV may differ on diagnosis, symptoms and condition, such as pregnancy and lactation. Understanding how these treatments change the vaginal microbiota are key in understanding how bacterial communities may change when exposed to antibiotics. Although BV cannot be attributed to a single bacteria, this shift to less protective species could influence its treatment, recovery and resurgence.

In a recent study, twenty-one pre-menopausal women with reports of abnormal vaginal discharge were recruited to discern how commonly prescribed treatments affected their vaginal microbiota [7]. The participants' treatment consisted of Metronidazole, which is the most commonly prescribed treatment for BV or Clindamycin, if the participant was pregnant or lactating

during the study. Vaginal swabs were taken during the first visit and at a four week follow up after treatment. With these samples, they used the Bray-Curtis dissimilarity to quantify how similar or different the samples in the cohort were [8]. Additionally, they used the Shannon index to determine the number of species and how these are distributed in the samples when compared to each other [8]. Through these analyses, two community types (CT) were identified, one composed of *Gardnerella*, *Prevotella*, *Atopobium* and *Sneathia* (CT 1), and the other had *Lactobacillus* dominance (CT 2). The patients were assigned one or the other based on their vaginal composition.

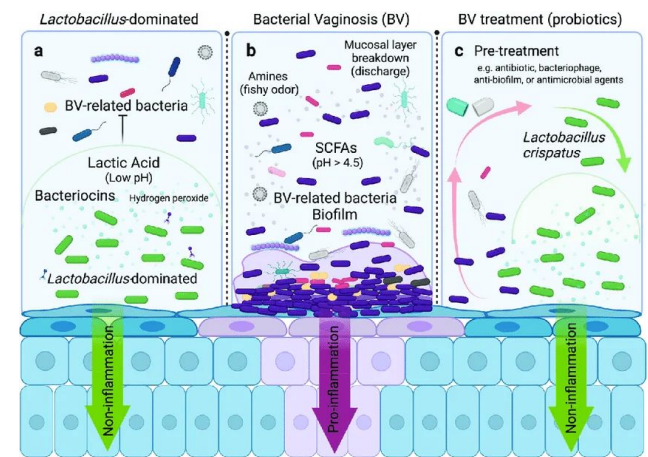


Figure 2: A visual comparison of the vaginal changes in a *Lactobacillus* dominated microbiota and one afflicted with BV. How treatment counteracts and is able to reestablish colonization of *Lactobacillus* [6].

After all antibiotic treatment, bacterial diversity decreased and the core microbiota consisted of *Lactobacillus*. The biggest change observed was in women who had a CT 1 composition, there was a significant decrease in the number of bacteria and distribution of these, as well as an increase in the presence of *Lactobacillus*, where previously they had been mainly colonized by anaerobes. Even after treatment, it is common that some women will experience recurrence. In this case, 43% of women experienced recurrence of abnormal vaginal, and 70% of women with persistent symptoms had been treated with Metronidazole.

The research team identified a possible association to a higher abundance of *Ureaplasma* in the participants who

experienced recurrence of symptoms. This is due to the fact that *Ureaplasma* is resistant to metronidazole but susceptible to clindamycin, and has been associated with BV recurrence. However, when considering the presence of *Ureaplasma* in women who experienced recurrence, no direct relationship was found between its presence and symptoms before treatment. Overall, this clinical study identified that the use of antibiotics does shift the microbiota to have a higher relative abundance of *Lactobacillus* [7].

Next Steps

Identifying a universal "healthy" microbiota has not been elucidated, as is the process of diagnosing BV and streamlining its treatment. Further understanding the vaginal microbiota and how it may shift to less protective populations, identifying possible universal markers for dysbiosis, as well as evaluating different treatments for BV may offer women alternatives that reduce recurrence.

By Saadia Jimenez-Ñeco

References:

- [1] V. S. Saraf et al., "*Vaginal microbiome: normalcy vs dysbiosis*," *Archives of Microbiology*, vol. 203, pp. 3793-3802, Jun. 2021.
- [2] J. Shen et al. "*Cervicovaginal microbiota: a promising direction for prevention and treatment in cervical cancer*," *Infectious Agents and Cancer*, vol. 19, no. 13, Apr. 2024.
- [3] L. A. Chacra, F. Fenollar and K. Diop, "*Bacterial Vaginosis: What Do We Currently Know?*," *Frontiers in Cellular and Infection Microbiology*, vol. 11, Jan. 2022.
- [4] K. Carlson and B. Mikes, "Amsel Criteria," *PubMed*, Jan 31, 2026.
- [5] D. Aleru, "Nugent Score," *Microbiome*, Mar 27, 2025.
- [6] S. Wu et al., "*The right bug in the right place: opportunities for bacterial vaginosis treatment*," *Nature Partner Journal Biofilms Microbes*, vol 8, no. 34, May 2022.
- [7] R. D. Zwittink, et al, "*The vaginal microbiota in the course of bacteria; vaginosis treatment*," *European Journal of Clinical Microbiology & Infectious Diseases*, vol 40, pp. 651-656, Oct. 2020.
- [8] I. Cassol, M. Ibañez and J. P. Bustamante, "*Key features and guidelines for the application of microbial alpha diversity metrics*," *Scientific Reports*, vol. 15, no. 622, Jan. 2025.



The Role of Epinephrine in Anaphylaxis

Severe allergic reactions are a widespread affliction. 51% of adults in the United States who have food allergies experience severe reactions, and about 1 to 5 out of 100 people in the United States are affected by anaphylaxis [1]. Epinephrine is a drug widely used to treat anaphylaxis or severe allergic reactions. It has significant and often lifesaving benefits, but also may have adverse consequences.

Anaphylaxis is a hypersensitivity disorder. In other words, it is an allergic reaction with an extreme degree of severity [2], [3]. Classified as an IgE-mediated hypersensitivity reaction, anaphylaxis is described as being acute, life-threatening, and multi-systemic [2]. While most allergic reactions typically only affect one area of the body, anaphylaxis may involve several organ systems [4].

An allergic reaction occurs when an innocuous substance (an allergen) is mistaken for a harmful one [3]. The most common allergens that cause anaphylaxis are foods such as peanuts, tree nuts, shellfish, and eggs, latex, insect stings, medications like penicillin and aspirin, anesthesia, and contrast dye [3]. This misconception induces the immune system to respond by producing a disproportionate amount immunoglobulin E, otherwise known

as IgE antibodies, which inundate the body with histamine [3], [5].

Histamine is a chemical that relays messages between cells, and it has a paramount role in allergic reactions. There are four types of histamine receptors; the resulting histamine-induced reaction depends on which receptor it binds to [4]. H1 receptors are prevalent all throughout the body and may incite negative effects. Some examples of possible reactions include flushing, pain, bronchoconstriction, vascular permeability, vasodilation, itching skin, tachycardia, and hypotension [4]. H2 receptors are primarily located within the heart, smooth muscle, and acid-releasing stomach cells. They may cause many of the same reactions as H1 receptors, as well as airway mucus gland stimulation and the secretion of stomach acid [4]. H3 receptors are generally situated in the central nervous system and play a role in the management of neurotransmitters. And finally, H4 receptors reside in hematopoietic cells and bone marrow, and are involved in inflammatory diseases and the development of particular blood cells [4].

There are several treatments for anaphylaxis depending on the symptoms, with antihistamines being commonly used to suppress further histamine effects. H1 blockers are used most commonly, while H2 blockers are used in extreme instances [2].

However, the main treatment for treating anaphylaxis is epinephrine. Epinephrine is a commonly used hormone and medication to treat several medical disorders, including anaphylaxis and other type 1 hypersensitivity reactions [6]. A more widely recognizable name for epinephrine is adrenaline. It is a central nervous system neurotransmitter that is secreted by the adrenal glands. This means that it chemically conveys messages between neurons and to cells.

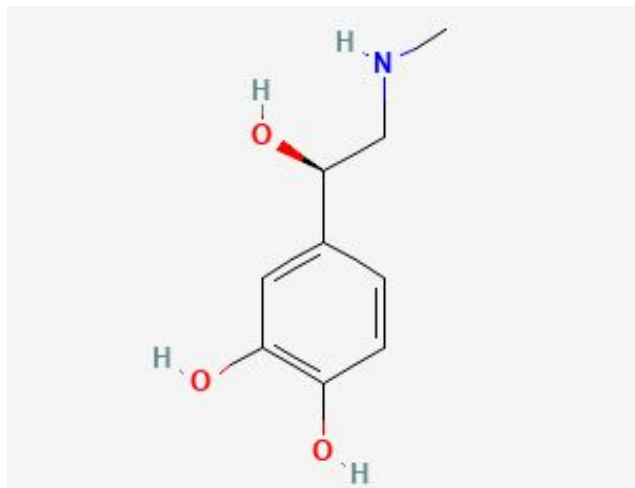


Fig 1: The molecular structure of epinephrine: $C_9H_{13}NO_3$ [7].

Epinephrine has several uses in the body as a neurotransmitter, including being involved in attention, metabolism, focus, excitement, and panic. However, epinephrine's main purpose, as both a hormone and a neurotransmitter, is its role in acute stress response, commonly known as "fight-or-flight" [8].

An acute stress response is the body's reaction to a stressful circumstance. During this response, the brain first discerns danger. This could be a physical or psychological peril. The brain then transmits a signal to the body via a neurotransmitter called noradrenaline, which then communicates to the adrenal glands. The adrenal glands excrete two hormones: epinephrine and noradrenaline. These hormones stay activated until the body leaves the perceived danger [8].

Physiological effects of epinephrine during this response include the skin paling, the pupils dilating, the heart rate speeding up,

the muscles becoming stronger from the increase of blood, the airways opening to receive more oxygen, and the liver releasing glucose [8]. Epinephrine acts on both alpha-adrenergic and beta-adrenergic receptors. Its action on these receptors makes it particularly useful for treating anaphylaxis. For example, it decreases vascular permeability, reduces vasodilation, and acts as a relaxant for bronchial smooth muscle [7].

Epinephrine treats anaphylaxis by improving breathing. It does so by relaxing airway muscles, constricting blood vessels, and decreasing swelling. It also impedes the chemicals which cause allergic reactions [6], [9]. Epinephrine should be administered as quickly as possible because the sooner it is given, the more likely it is to have a favorable outcome; more treatment becomes necessary as the reaction gains severity over time [10].

Epinephrine has several uses as a medicine besides treating anaphylaxis, including being used for cardiac arrest, asthma, and septic shock [6]. However, treatment does not ensure recovery. Sometimes, the body will experience a second anaphylactic reaction. This is called biphasic anaphylaxis. Biphasic anaphylaxis may occur even if the first treatment was successful and there has been no further exposure to the allergen. It is estimated to transpire in up to 20% of anaphylaxis cases. While biphasic anaphylaxis is most likely to happen 8 to 11 hours after the initial anaphylactic reaction, it may occur up to 72 hours later. Symptoms of biphasic anaphylactic reactions tend to be more mild than the first reaction, but this is not always the case. The most common treatment of biphasic anaphylaxis is also epinephrine [11].

Epinephrine can be administered for anaphylaxis in several different ways. It may be inhaled with a nasal spray or injected with a standard syringe. An under the tongue, or sublingual, film may also be used, or epinephrine may be auto-injected with a self-injectable syringe. One of the most widely known devices for epinephrine

administration, the EpiPen, is an example of an auto-injector [12].

While epinephrine has several life-saving benefits, it also may cause various side effects in miscellaneous areas of the body. This includes side effects which are related to the central nervous system, the cardiovascular system, the dermatologic system, the endocrine system, the gastrointestinal system, the neuromuscular system, the renal system, and the respiratory system. Some specific examples include agitation, nervousness, anxiety, nausea, vomiting, weakness, and tremors [6].

Drug-drug interactions may also pose complications. Drugs such as vasodilators, α -blockers, or antihypertensives may reduce or block the vasoconstrictive effects of epinephrine. This could counteract epinephrine's ability to increase blood pressure. Other drugs, such as β -blockers and antidepressants, may augment the effects of epinephrine. These are only a few examples of several drug-drug interactions which may disrupt the effects of epinephrine [6].

Other cons of epinephrine include that it can

become expensive due to the fact that it is not particularly long-lasting, so the treatment may need to be repeated. It is also expirable, meaning it must be replaced after time, even if it goes unused. Furthermore, epinephrine can be difficult to store because it must be kept at room temperature. This can be challenging for individuals who reside in areas that are exposed to extreme temperatures [13].

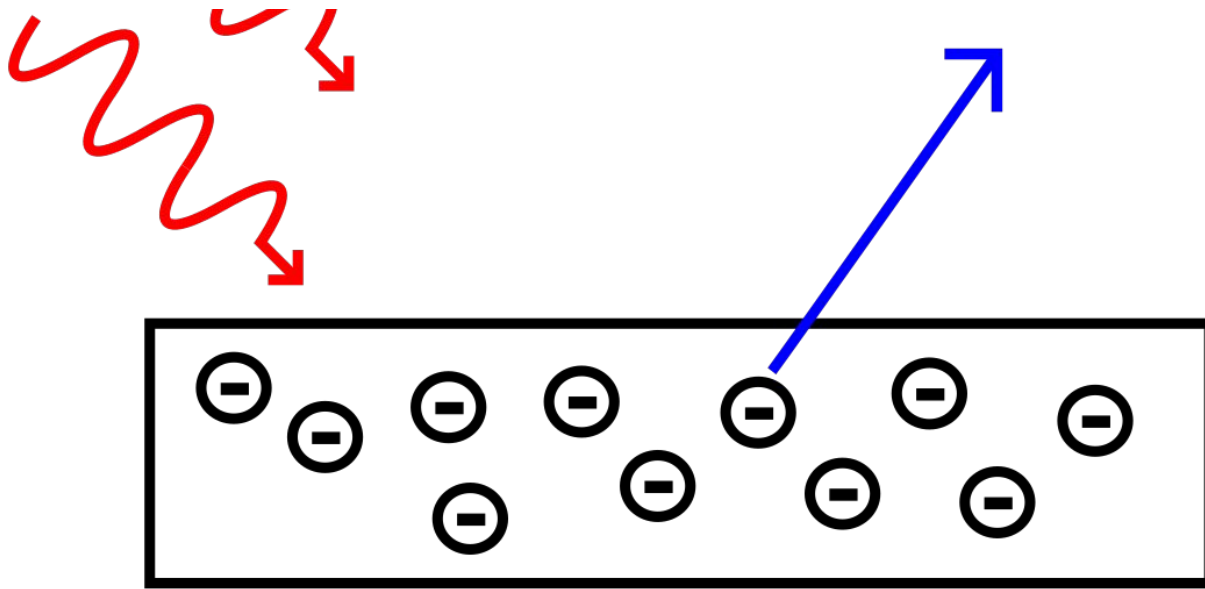
However, in many scenarios, epinephrine's benefits outweigh its drawbacks. Besides its strong lifesaving-potential for anaphylaxis, its availability in self-injectable syringes makes it easy to administer in emergencies, even for medically inexperienced individuals. It is also suitable to be used on pregnant or breastfeeding women [12], making it usable for a larger and less restricted portion of the population.

Epinephrine is a versatile and life-saving treatment which is used to treat anaphylaxis. While the treatment has several drawbacks, they are often outweighed by its benefits. Anaphylaxis would cause more harm without the aid of epinephrine than it does with its treatment.

By Sara Glaesser

References:

- [1] "Anaphylaxis Statistics," Allergy & Asthma Network.
- [2] K. McLendon and B. T. Sternard, "Anaphylaxis," National Library of Medicine, 2023.
- [3] M. Armitage, "What is Anaphylaxis? Symptoms, Causes, Treatment and More," GoodRx, Aug. 25, 2021.
- [4] "Histamine," Cleveland Clinic, Mar. 28, 2023.
- [5] "Anaphylaxis," Allergy & Asthma Network, 2024.
- [6] R. Dalal and D. Grujic, "Epinephrine," National Library of Medicine, 2024.
- [7] "Epinephrine," NIH.gov, 2019.
- [8] "Epinephrine (Adrenaline)," Cleveland Clinic, 2022.
- [9] R. Tamayo, "Epinephrine as a Treatment for Anaphylaxis," Penn Medicine Becker ENT & Allergy, Jul. 28, 2023.
- [10] G. Cloves, "All About Epinephrine: What it Does in a Reaction, How Long it Lasts, When it Gets Hot or Cold," Living Allergic, 2026.
- [11] S. Morris, "Everything You Should Know About Biphasic Anaphylaxis," Healthline, Feb. 20, 2026.
- [12] "Routes of Administration of Epinephrine for Anaphylaxis," Allergy & Asthma Network, Dec. 30, 2024.
- [13] R. C. Hakim and K. Hannemann, "Epinephrine (EpiPen): Uses, Alternatives, Side Effects & More," GoodRx, Mar. 03, 2022.



Photoelectric Effect And Its Practical Role in AI And Hardware

Introduction

For decades the semiconductor industry has been sustaining Moore's law, which states that the number of transistors on a microchip doubles every two years [1]. However, as transistors continued shrinking, approaching atomic sizes, the industry faced a problem. This scale makes electrons prone to quantum tunnelling, which is a quantum mechanical phenomenon where an electron passes through a physical barrier that would be impassable in classical physics. This causes leakage currents, leading to power inefficiency and errors, making further miniaturisation impractical [2].

One possible solution to maintain computer performance improvements is shifting from purely electron-based computing to implementing the photoelectric effect in hardware. The photoelectric effect was first described by Albert Einstein in 1905, later earning him a Nobel Prize in 1921. Einstein theorised that light exhibits wave-particle duality, travelling simultaneously as a continuous wave and as discrete energy packets called photons. In the photoelectric effect, a photon of sufficient energy is absorbed by an atom, causing it to eject an electron from its orbital [3].

Photoelectric Effect

The photoelectric effect is described by a formula, $E = h\nu$, where h is Planck's constant (6.626×10^{-34} Js) and ν is the threshold frequency at which the photoelectric effect occurs [3]. When $h\nu$ exceeds the material's work function Φ , electron emission occurs, with excess energy becoming the electrons' kinetic energy ($E_k = h\nu - \Phi$) [4]. This behaviour makes the photoelectric effect useful as an engineering tool. A photodetector material is designed with a specific work function matched to an incoming photon wavelength; it responds only to that specific wavelength.

Modern silicon photodetectors operate at a standard telecommunication wavelength of 1550 nm, and they achieve quantum efficiencies of around 90%. Consequently, when a photon hits the silicon detector, 9 in 10 incoming photons successfully knock an electron loose, leaving behind a positively charged hole. This electron-hole pair generates a measurable electrical signal [5]. The vacancy in the innermost orbital of the atom, the K-shell, that is left after the photoemission triggers a cascade as an outer-shell electron drops to fill it, releasing the energy difference either as another electron (the Auger electron) or an X-ray. The probability of a photon hitting an electron increases proportionately to the atomic number to the power of 4.5 [6]. As a result, lead is widely used in radiation shielding applications due to its high atomic

number (82) [7].

In production, silicon microchips are most common due to their low costs of production, yet they do not obtain the greatest quantum efficiencies. To balance the cost of production and improve probabilities of photons hitting electrons, engineers use silicon along the waveguides and germanium at the end of the circuits. Since germanium, with an atomic number of 32, provides a narrower bandgap.

Photonics Circuits

In a direct bandgap semiconductor, made out of materials such as gallium arsenide (GaAs) or indium phosphide (InP), the lowest energy state of the conduction band lines up perfectly with the highest energy state of the valence band in momentum space. Consequently, when an electron drops down, it emits excess energy in the form of a photon that can carry data. However, in an indirect bandgap semiconductor commonly made out of silicon (Si), when an electron drops down, it has to change its momentum, releasing heat [8], [9].

This type of semiconductor is used to let photons flow, yet they have a risk of overheating; to minimise that risk, engineers use heterogeneous integration of indium phosphide laser sources, silicon substrates and an emerging thin film made out of lithium niobate on insulator (LNOI).

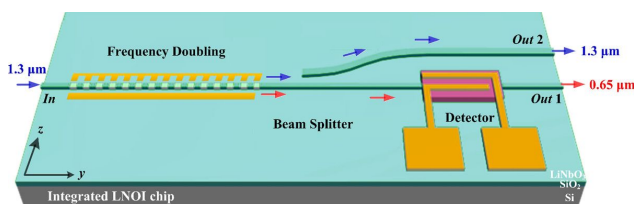


Figure 1: Diagram of the lithium niobate on insulator [10].

In this heterogeneous integration at the emitter, data is encoded onto a continuous-wave laser using a modulator, typically a Mach-Zehnder interferometer, that encodes binary information by shifting the phase of the optical arm relative to the other, producing constructive or destructive interference at the output [11], [12].

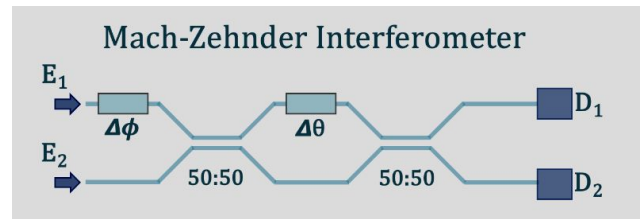


Figure 2: Schematic diagram of an ideal Mach-Zehnder interferometer with perfect split ratios [13].

Those encoded photons travel through silicon waveguides that are 500 nm wide channels fabricated on a silicon-on-insulator substrate. Inside a silicon core, light slows to 85600 km/s, because the surrounding silicon dioxide allows the light to travel faster; the refractive index contrast is large enough to bend a waveguide, consequently guiding the light with propagation losses around 2-3 dB/cm in standard SOI processes [14]. Because photons carry no electrical charge, they do not create any electrical resistance with the silicon dioxide; consequently, there is no Joule heating. Instead, we can encounter other issues such as a transmission loss from scattering light on sidewall roughness and photon absorption by free carriers [15].

Moving forward, wavelength division multiplexing exploits the non-interacting nature of photons, allowing multiple independent data streams at different frequencies to propagate simultaneously through the same waveguide; at the receiver, they are being demultiplexed by the ring resonators or arrayed waveguide gratings tuned to specific channels. However, scaling Wavelength Division Multiplexing systems to numerous discrete lasers in one waveguide causes a risk of thermal overheating of the chip. To bypass this risk, we use the Kerr frequency comb, a light source that creates multiple equally spaced optical channels. By feeding a single pump laser into a silicon microresonator, dozens of evenly spaced wavelength channels are being generated through nonlinear four-wave mixing, which now enables dense wavelength division multiplexing without arrays of discrete lasers [16].

At the destination, a photodetector converts

incoming photons back into electricity. Incoming photons are absorbed by the detector material, knocking electrons loose through the photoelectric effect. The photocurrent that is being produced by this is the electrical reconstruction of the original binary signal.

Josephson Junctions

Josephson junctions are circuits based on quantum tunnelling; they consist of two superconducting electrodes separated by an insulating barrier of only 1-2 nm [17]. Below the superconducting critical temperature, electrons form Cooper pairs, which are bound pairs of electrons at low temperatures that tunnel through the barrier, producing a supercurrent with zero applied voltage. When voltage is applied, the tunnelling current oscillates at a frequency $f = 2eV/h$ with a proportional constant which defines the SI volt. Photon-assisted tunnelling connects to photonic systems, where photons control tunnelling below the classical threshold via discrete energy quanta. As a switching element, a Josephson junction dissipates approximately 10^{-23} joules per event [18].

Niobium-based junctions require temperatures below 9.2K to operate, and quantum computing pushes for temperatures below 1K using multi-layered dilution refrigerators, made out of copper and gold, which makes them resemble golden chandeliers. Yet even at such low temperatures, decoherence driven by residual thermal fluctuations and material defects limits qubit coherence times to microseconds [19]. The refrigeration infrastructure consumes kilowatts to maintain millikelvin temperatures, and that overhead dominates the energy budget entirely; the computational efficiency gains are largely cancelled out by the cost of maintenance [20]. For AI data centres lacking specialised cooling infrastructure and operating near room temperature, quantum computing will be impossible to

deploy with minimal decoherence.

Current Photonic Hardware

While quantum computing remains solely experimental [21], photonic systems are already used in some of big tech's hardware. In the Dell PowerSwitch SN6000 [22], NVIDIA Spectrum-X Ethernet Photonics delivers co-packaged optical (CPO) technology. This technology integrates silicon photonics in the same package as an application-specific integrated circuit (ASIC), minimising copper traces [23]. Intel's integrated optical compute interconnect (OCI) chiplet co-packaged with an Intel CPU supports 64 channels up to 100 metres of fibre optics. That is expected to address AI's infrastructure's growing demand for higher bandwidth, lower power consumption and longer reach [24].

Lightmatter's Envisi is a hybrid photonic-electronic processor, integrating 6 chips combining 50 billion transistors and 1 million photonic components. Envisi executes accelerated matrix-vector computations, using Mach-Zehnder interferometers for light wave interactions. Through wavelength division multiplexing, it drastically increases computational density [25]. The integration of these photonic systems into hardware fixes the physical limits of copper wiring in AI chips [26]. However, silicon waveguides remain sensitive to temperature changes [14], which requires hardware to integrate on-chip microheaters [27].

Conclusion

The photoelectric effect serves as the foundation for this advanced computing infrastructure. Photonic engineering directly addresses the physical limit where copper wires can no longer scale down to match shrinking transistors. Although thermal stability remains a challenge, it is strictly a practical engineering hurdle rather than a fundamental physical limit

By Weronika Padjasek

References:

- [1] G. E. Moore, "Cramming more components onto integrated circuits," *Electronics*, vol. 38, no. 8, pp. 114–117, Apr. 1965.
- [2] IEEE IRDS, "More Moore," 2022 Edition.
- [3] A. Einstein, "Über einen die Erzeugung und Verwandlung des Lichts betreffenden heuristischen Gesichtspunkt," (in German), *Ann. Phys.*, vol. 17, pp. 132–148, 1905.
- [4] "Neutral Particle Detector," IPPP, Durham Univ.
- [5] S. M. Sze and K. K. Ng, "Physics of Semiconductor Devices," 3rd ed. Hoboken, NJ, USA: Wiley, 2006.
- [6] B. K. Agarwal, "X-Ray Spectroscopy," 2nd ed. Berlin, Germany: Springer-Verlag, 1991.
- [7] E. M. Attia et al., "Mechanical and radiation shielding assessment of high-density/magnetite-cementitious plaster for repair and retrofit of nuclear power plant structures," Jan. 10, 2026.
- [8] T. Kita et al., "Fundamentals of semiconductors," Springer, ch. 8.
- [9] H. Ueki, "Various semiconductors," QPS Lab, 2024.
- [10] J. Heei et al., "Heterogeneously integrated lithium niobate on insulator chip for frequency up-conversion and detection," *ACS Photon.*, 2024.
- [11] B. Jalali and S. Fathpour, "Silicon photonics," *J. Lightw. Technol.*, vol. 24, no. 12, pp. 4600–4615, Dec. 2006.
- [12] A. R. Barbosa, F. Carlesso, Y. Abe, M. I. Alayo, and L. E. A. Vieira, "Mach-Zehnder interferometers in photonic circuits for solar irradiance investigations," in *Proc. Seminatec*, 2024.
- [13] C. G. Valdez et al., "High contrast nulling in photonic meshes through architectural redundancy," 2025.
- [14] S. Talabattula, "Polarization analysis of an asymmetrically etched rib waveguide coupler for sensing applications," *J. Lightw. Technol.*, vol. 23, no. 12, pp. 4222–4238, Dec. 2005.
- [15] J. Zhang et al., "New-generation silicon photonics beyond the singlemode regime," 2021.
- [16] T. J. Kippenberg et al., "Dissipative Kerr solitons in optical microresonators," *Science*, vol. 361, no. 6402, Aug. 2018.
- [17] B. D. Josephson, "Possible new effects in superconductive tunnelling," *Phys. Lett.*, vol. 1, no. 7, pp. 251–253, 1962.
- [18] M. H. Devoret and R. J. Schoelkopf, "Superconducting circuits for quantum information: An outlook," *Science*, vol. 339, no. 6124, pp. 1169–1174, 2013.
- [19] C. E. Murray et al., "Decoherence in superconducting qubits," *Phys. Rev. Res.*, vol. 4, no. 4, 2022, Art. no. L042038.
- [20] M. Fellous-Asiani, J. H. Quilter, and P. Campagne-Iba, "Energy overhead of superconducting quantum computing," 2021.
- [21] Plain Concepts, "Quantum computing potential challenges."
- [22] BlueAlly, "PowerSwitch SN6800 LD Spec Sheet."
- [23] Y. Liu et al., "Co-packaged optics for data centers," *Front. Optoelectron.*, 2022.
- [24] Intel Newsroom, "Intel unveils first integrated optical I/O chiplet," Jun. 26, 2024.
- [25] Lightmatter Blog, "A new kind of computer," 2025.
- [26] HPCwire, "Lightmatter introduces rack-on-chip interconnect to increase performance and reduce energy."
- [27] F. Gan et al., "Maximizing the thermo-optic tuning range of silicon photonic structures," in *Proc. Photon. Switching*, pp. 67–68, 2007.

Thank You!

This concludes the ninth issue of the Penrose Magazine. Thank you so much for reading and we hope you enjoyed!

Finally, we would like to thank our team for all the work they put into editing, promoting, and formatting this issue.

If you would like to write for the next issue of the Penrose Magazine, please check out our website or socials:

[Penrosemagazine.org](https://penrosemagazine.org) | [Instagram](#) | [TikTok](#)

