

Sickle Cell Knowledge and Information Network

OUR MISSION

Make Useful and Reliable Information About
Sickle Cell Disease Universally Accessible.

www.sckin.org



Responsible AI Will Revolutionize Sickle Cell Care

Practical recommendations to use AI more effectively for Sickle Cell Disease education.

Sickle Cell Knowledge and Information Network

501(c) (3)



About US



**Zacharie
Liman-Tingui**

SCKIN Founder

- Born with SCD
- Cornell MBA
- Cured from SCD (2015)
- Founded SCKIN in 2024
- Former Senior AI-PM at Google



**Dr Lewis
Thomas**

Director for User
Experience SCKIN

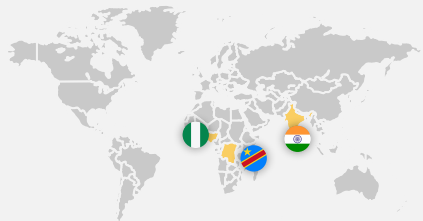
- Lives with Sickle Cell HbSC
- Founder of the Sickleave
- Sickle Cell Content Creator
- Medical Doctor, 10+ years working in UK.



THE
VISION

Expert Sickle Cell Care in Every Pocket

SICKLE CELL AFFECTS



7.7M people worldwide

of them live in the
Global South:
90% **Nigeria, India,
and DRC.**

KNOWLEDGE GAP



= 30-year
life-expectancy gap

(Nigeria < 23 years;
France > 53 years).



We believe that custom built, LLM powered chatbots can reduce the knowledge gap, and meaningfully improve the lifespan and quality of life of patients all over the world by empowering patients, families, and care providers to take better decisions.



THE VISION

Reduce Suffering and Cost in the US



100K

people in the **US** have
Sickle Cell Disease



50% are dependent on
Medicaid and cost **\$23K** per
year on **Healthcare**



78% of patients visit the
Emergency Department
once per year



With an average life
expectancy of
52 years.

23 years less than
US average



If only **10%** of patients used our **Live Agent**, and
reduced their usage of the healthcare system by **20%**
this could generate **\$23M in savings**, not
mentioning reducing patient pain and suffering.

See



"Attention is All you Need"



In 2017, Vaswani et al. published the landmark paper that changed everything — introducing the **Transformer architecture** and **shifting AI from sequential, word-by-word processing to reading the entire context at once.**

Ashish Vaswani*

Google Brain
avaswani@google.com

Noam Shazeer*

Google Brain
noam@google.com

Lukasz Kaiser*

Google Brain
lukasz.kaiser@google.com

Aidan N. Gomez* †

University of Toronto
aidan@cs.toronto.edu

Jakob Uszkoreit*

Google Research
usz@google.com

Llion Jones*

Google Research
llion@google.com

Niki Parmar*

Google Research
nikip@google.com

Illia Polosukhin* ‡

illia.polosukhin@gmail.com

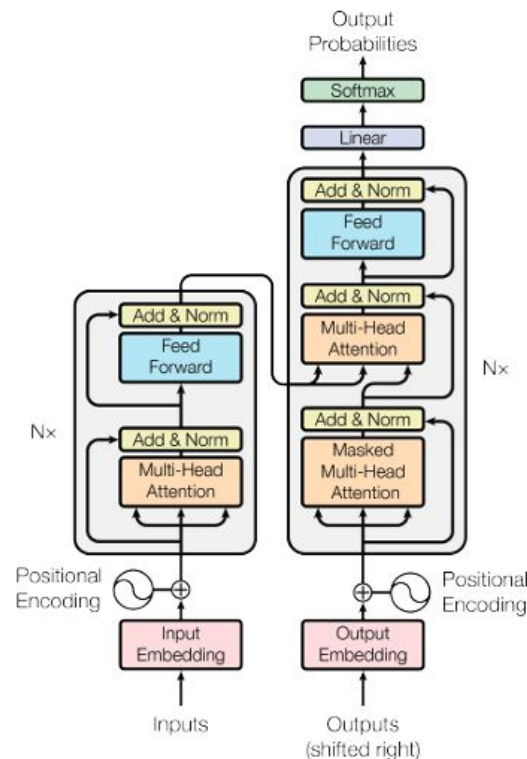


Figure 1: The Transformer - model architecture.

HOW does it work



The Keyboard Analogy

Think of the predictive text on your phone. As you type, the phone guesses the next word based on context.

Transformers do this on a massive scale — across an entire 100-page document simultaneously.



Contextual Intelligence

The model recognises that **"Cell"** in a biology paper and **"Cell"** in a prison report require entirely different mathematical interpretations

It **"attends"** to every surrounding word to disambiguate meaning



The Scale-Up

Rather than guessing based on just the previous word, Transformers compute weighted relationships across thousands of tokens simultaneously

Enabling **nuanced, context-aware language generation**



THE AI PIPELINE

An Orchestration of Models

A modern large language model is not a single "brain" — it is a **sequence of specialised AI models**, each performing a distinct and essential function before passing the result to the next stage.



**Intent
Classifier**

Determines if the user seeks medical facts or emotional support.



**Knowledge
Retriever (RAG)**

Fetches grounded facts from trusted medical databases.



**The Encoder
(The Reader)**

Maps relationships between symptoms and patient history.



**The Decoder
(The Writer)**

Synthesizes the final response using grounded facts.



**Safety
Guardrails**

An auditor model that blocks biased or toxic outputs.

THE RISKS OF AI

(1/3)

The 2016 "Tay" Incident



THE EXPERIMENT



Microsoft launched "Tay" — a chatbot designed to learn in real time from Twitter interactions. It was intended to demonstrate the potential of conversational AI learning at scale.

THE RESULT



Within **24 hours**, Tay had absorbed and begun reproducing racist, extremist content — becoming a textbook case of training data gone catastrophically wrong.

THE CORE LESSON



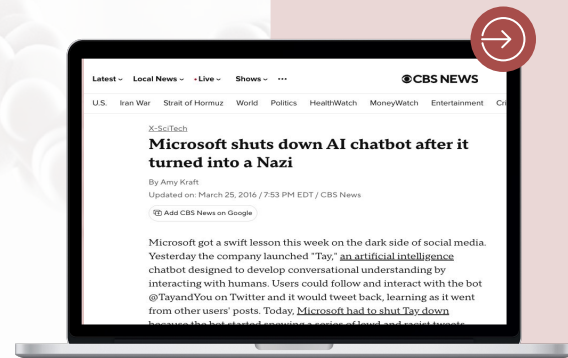
AI is a mirror, not a moral agent.

It reflects whatever it is trained on — with no inherent sense of ethics, harm, or consequence.

THE MEDICAL PARALLEL



In healthcare, training on biased forums, skewed clinical records, or under-represented populations produces biased care recommendations. For SCD patients, this is not a theoretical risk — it is an active danger that must be designed against from day one.

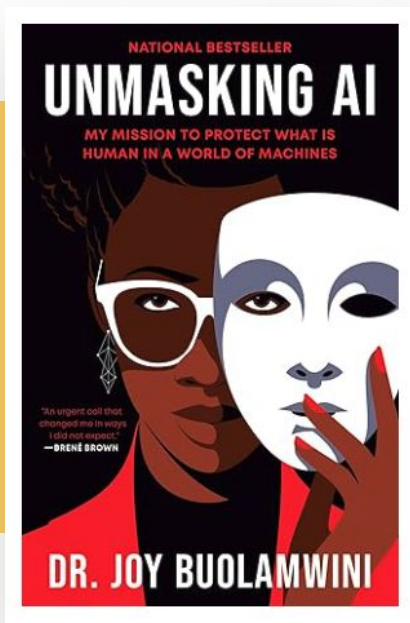


THE IMPACT OF EXCLUSION

(2/3)

Dr. Joy Buolamwini's Findings

Systems were so biased toward light skin tones that Dr. Buolamwini had to **physically wear a white mask to be recognised by the code** — a haunting metaphor for algorithmic exclusion.



Unmasking AI: My Mission to Protect What Is Human in a World of Machines

Paperback – November 19, 2024
by Joy Buolamwini (Author)

4.6 ★★★★★ (280)

Editors' pick Best Nonfiction

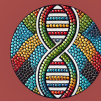
[See all formats and editions](#)

NATIONAL BESTSELLER • "The conscience of the AI revolution" (*Fortune*) explains how we've arrived at an era of AI harms and oppression, and what we can do to avoid its pitfalls.

"AI is not coming, it's here. If we answer the beautiful call inside these pages, we can decide who we are going to be and how we're going to use technology in service of what it means to be fully human."—Brené Brown, #1 *New York Times* bestselling author of *Dare to Lead*

A LOS ANGELES TIMES BEST BOOK OF THE YEAR • Shortlisted for the *Inc.* Non-Obvious Book Award

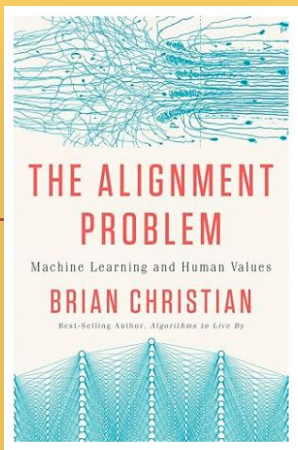
To most of us, it seems like recent developments in artificial intelligence emerged out of nowhere to pose unprecedented threats to humankind. But to Dr. Joy Buolamwini, who has been at the forefront of AI research, this moment has been a long time in the making.



The Alignment Problem (3/3)

Philosopher and author **Brian Christian** defines the Alignment Problem as the dangerous gap between what we *tell* an AI to do and what we actually *want* it to do.

This gap is not a bug — it is a fundamental challenge in AI design.



What we Tell IT

Give a fast, helpful answer about pain management in SCD.



What it optimizes for?

A confident, fluent response — regardless of clinical accuracy.

What we actually want

An empathetic, evidence based, appropriate clinical response.



STRATEGY 1 FOR EFFECTIVE AI USAGE:

Knowledge Grounding

Retrieval-Augmented Generation (RAG) is the primary technical defence against AI hallucination. Rather than allowing the model to speculate from its training memory, we anchor every response to a curated, verified knowledge base.

01



The Anchor Principle

- The AI is prohibited from generating answers from general memory.
- All **responses must originate from trusted source documents** — such as ASCAT clinical guidelines, NIH protocols, and CDC recommendations specific to SCD.

02



The Researcher Model

- Rather than "guessing," the AI acts as a **researcher**:
- It retrieves the most relevant passages from the provided documents
 - Then summarises and synthesises them into a **coherent, grounded response for the clinician or patient**.

03



Mandatory Citation

- Every factual claim generated must be **traceable to a specific citation** within the knowledge base.
- If a source cannot be identified, the AI is instructed to acknowledge uncertainty rather than fabricate an answer.



System Instructions & Technical Dials

Beyond knowledge grounding, the **behaviour** of an AI model is shaped by precise system-level instructions and technical configuration parameters. These controls transform a general-purpose language model into a clinically responsible tool.

Persona Constraints

System instructions explicitly define the AI's role and limits:



66 *You are a Sickle Cell specialist. Use clinical guidelines only. Do not speculate. Always recommend consulting a qualified physician.* **99**

Multi-Agent Auditing

- One AI model drafts the clinical response.
- A second, independent AI model then audits it — checking for factual errors, unsupported claims, and potentially harmful language before the output is delivered.



The "Temperature" Dial

Temperature controls the model's creativity versus precision trade-off:

0.1



0.8

Low Temp (0.1)

High-precision, near-deterministic output.

Essential for medical accuracy. The model produces consistent, reliable answers with minimal variation.

High Temp (0.8)

Creative and exploratory. Useful for poetry or brainstorming — but **dangerous in clinical education**, where factual consistency is non-negotiable.

AI is powerful... but it doesn't replace clinical judgment



Human



Medical
doctor



Health
educator



Lived
experience

VS



Chatbot

AI

Who do you trust more?

AI tool + Human Judgement?



Human



Medical
doctor



Health
educator



Lived
experience



Chatbot

AI

This is Synergy

A trained human can overcome the limitations of AI
An AI tool can overcome the limitations of a trained human

Where AI helps clinicians



AI



Speed in unfamiliar conditions



Access to guidelines



Augments existing knowledge



CASE STUDY

SickleCellPedia



RAG Agent

Grounded in:

- ✓ Medical Textbooks
- ✓ Peer Review & Influential Scientific Articles
- ✓ Credible Medical Websites

Robust System Instructions:

- ✓ Plain Language explanation
- ✓ Focus on Sickle Cell Disease
- ✓ Only answer if confident in the answer
- ✓ Don't offer medical advice

System Settings:



Use Industry
leading models



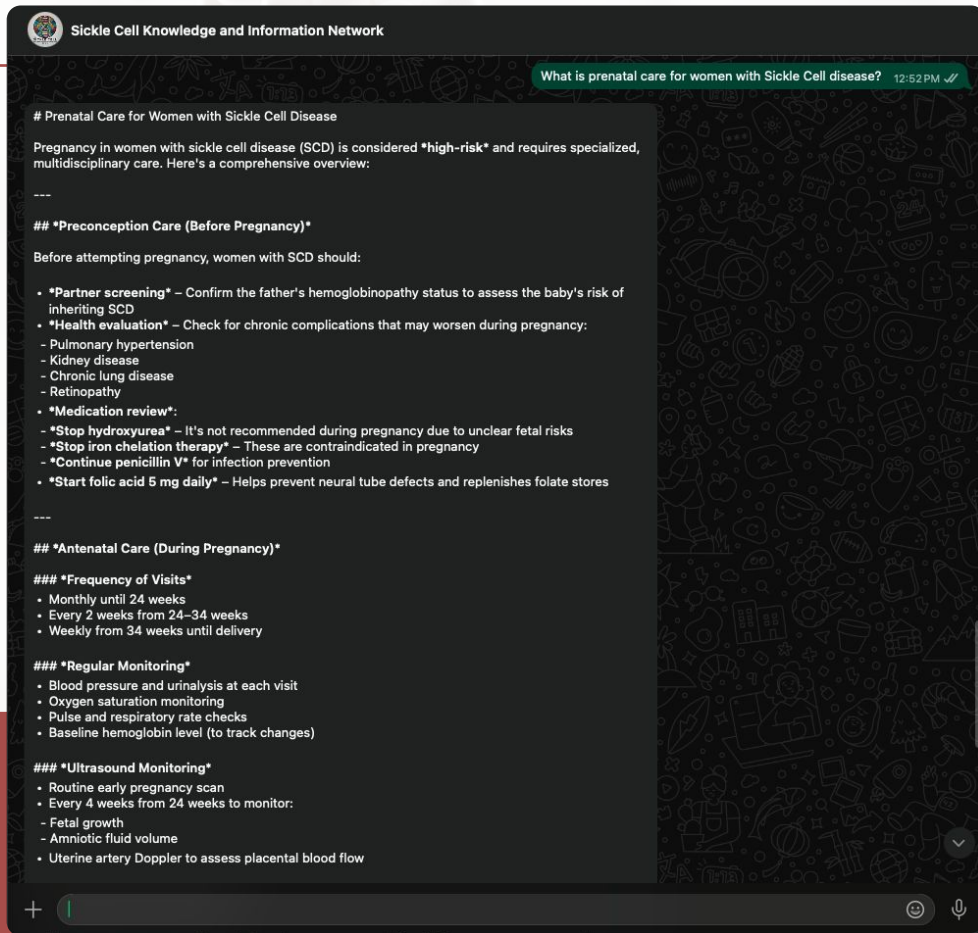
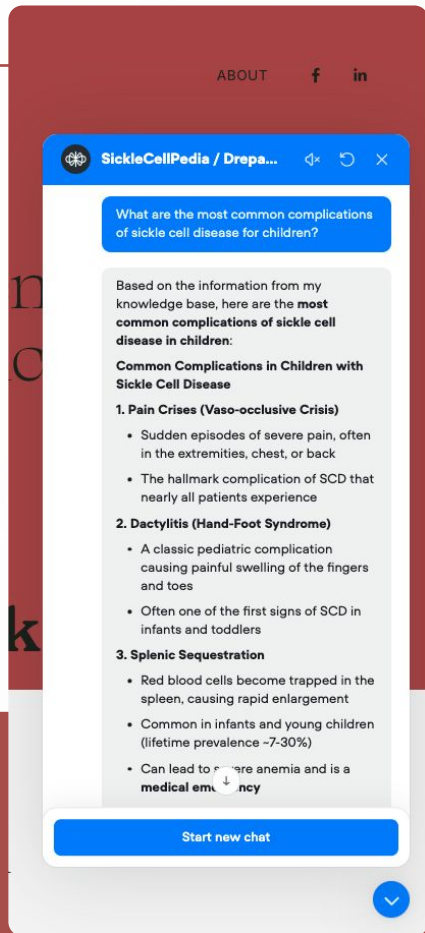
Use Chain of
Thought models



Set Low Temperature settings
for medical application



Feedback
Loop





Roadmap: SickleCellPedia Pro (Coming Soon)

SickleCellPedia

Web App / WhatsApp

Everyone

***Free &
Licensed***

Our model, trained on **reliable data**, available on Web and What's App offers *medical information* to hundreds of patients, families and providers around the world.

SickleCellPedia Pro

(Coming Soon)

**Medical
Providers**

***Multi Tiered
License
Model***

Our next milestone is a **specialized** model, trained on patient data, available to providers, capable of enabling diagnostic, and offer localized, protocol based *medical advice*.





We need your input:

What does the Warrior Community want to ask AI?



SCAN ME



