

2026

AI JOURNEYS AT MICHIGAN

RESEARCH THROUGH THE AI LIFECYCLE

TABLE OF CONTENTS

2	INTRODUCTION
3	PHASE 1 RESEARCH DESIGN
5	PHASE 2 DATA ACCESS AND DISCOVERY
7	PHASE 3 DATA PREPARATION
9	PHASE 4 COMPUTE RESOURCES
11	PHASE 5 BASELINE ANALYSIS
13	PHASE 6 ADVANCED ANALYSIS
15	PHASE 7 CUSTOM ENGINEERING
17	PHASE 8 OUTPUTS AND REPRODUCIBILITY
19	PHASE 9 KNOWLEDGE SHARING
21	CONCLUSION

INTRODUCTION

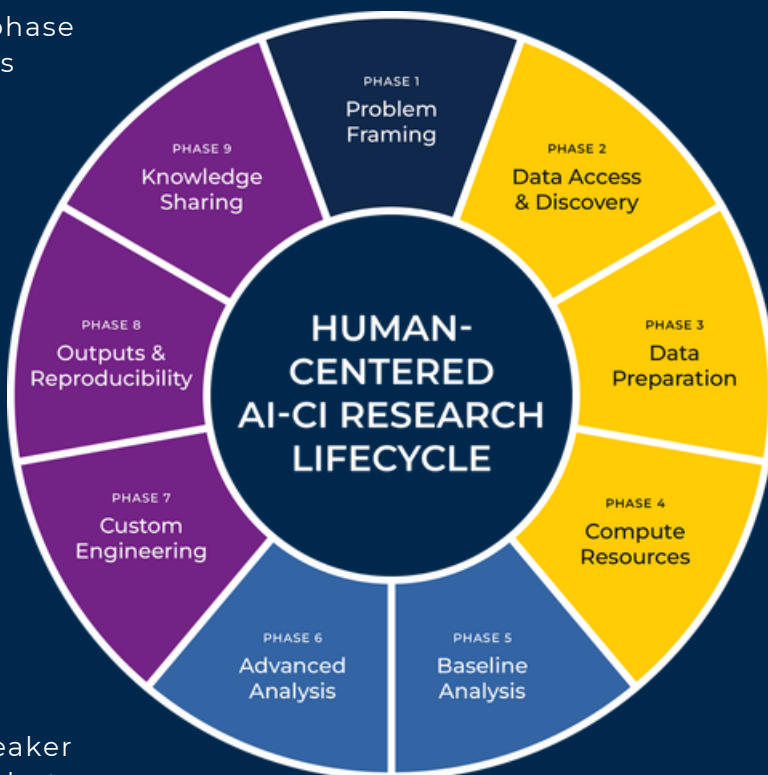
THE AI RESEARCH LIFECYCLE

Artificial intelligence is no longer a single tool dropped into the middle of a research project — it has become woven through every phase of the scientific process, from the first question a researcher asks to the moment findings are shared with the world. Yet the pace of AI adoption has far outrun the frameworks researchers have for thinking about it clearly.

The AI Research Lifecycle is a nine-phase model developed to help researchers navigate this landscape with intention. It does not prescribe a single workflow. Instead, it names the key decision points at which AI can either amplify a researcher's contributions or quietly introduce assumptions, biases, and blind spots that undermine them. Understanding which phase you are in — and what AI can and cannot do for you there — is a form of scientific literacy for the twenty-first century.

The talks collected in this booklet were delivered by University of Michigan researchers as part of MIDAS's AI Journeys series. Each speaker shared not a polished success story, but an honest account of how they met AI in their own work: the dead ends, the pivots, the unexpected discoveries. Taken together, they map onto almost every phase of the lifecycle, offering a rare ground-level view of what AI-enabled research actually looks like across disciplines — from neurosurgery to power engineering, computational chemistry to social work, musicology to child welfare.

The nine phases are presented sequentially in this booklet, each introduced briefly and then illustrated by the researchers whose journeys most vividly embody it. Not every phase is represented by a single talk; some speakers move through several. And the lifecycle itself is not strictly linear — researchers cycle back, skip ahead, and discover that a problem in Phase 7 reaches all the way back to a choice made in Phase 1. That messiness is part of the point.



To learn more about each individual AI Journey, visit our website at midas.umich.edu/2026-journeys or watch their presentation on our YouTube channel ([myumi.ch/393pz](https://www.youtube.com/channel/UCmyumi.ch/393pz)).

Framing the Question Before Reaching for the Tool

The first and arguably most consequential phase of any AI-enabled research project is not writing code or training a model — it is deciding whether, how, and why AI should be involved at all. Researchers at this stage must assess what assumptions AI introduces, where theory and domain expertise remain irreplaceable, and whether their question is genuinely enabled by AI or merely dressed up in it. Getting Phase 1 right does not guarantee a successful project; getting it wrong almost certainly guarantees problems that no amount of computational power can fix downstream.



Ali Namvar

Galban Lab — STREAM ICU monitoring framework

Namvar's development of STREAM, an explainable AI framework for real-time ICU patient monitoring, is a study in following the question rather than the tool. Faced with a high-dimensional, continuously streaming data environment, he identified that existing systems collapsed patient complexity into a single score or fired alerts without directional context. The choice of optimal transport — a method that measures the minimum work required to reshape one probability distribution into another — was not made because it was fashionable, but because the clinical question demanded a method that could track how an entire patient population shifted across physiological states over time. "We follow the question, not the tools," he said. "Optimal transport was a natural fit." This methodological honesty also shaped his team composition: he recognised that no single discipline could frame, validate, and interpret the outputs, bringing together engineers, clinicians, and biologists from the outset.



Jingyi Qiu

School of Information — measuring AI hype in scientific papers

Qiu's journey began not with a dataset but with a nagging observation: that paper abstracts were routinely overstating what the results actually showed, and that this gap was widening with the spread of large language models. The research design challenge was acute — how do you measure rhetorical inflation objectively, at scale, without a model that itself introduces the biases you are trying to measure? Her first attempt, asking an LLM to rate hype directly, failed because LLMs lacked the domain expertise and produced coarse, collapsed scores. Her second, generating a continuous rhetorical spectrum via chain-based sampling, failed because LLMs anchor too strongly on reference examples. Only on her third attempt — using persona-controlled writing — did the design work. Each failure was a Phase 1 reckoning with what AI can and cannot reliably do as a measurement instrument.



Farnaz Jahanbakhsh

CSE & School of Information — personalised content moderation

Before a single line of her content-moderation system was written, Jahanbakhsh conducted a formative qualitative study with people who had phobias, PTSD, and other sensitivities — not to gather training data, but to understand the structure of the problem. The research design insight that followed was that harm is personal, temporal, and contextual, and that no universal AI-based suppression policy could capture it. That finding determined everything: the shift from platform-level moderation to recipient-side personalised transformation, the decision to let users define their own sensitivities in natural language, and the ethical constraint that the system must never silence the original poster. AI was brought in only once the problem was understood well enough to know what it should and should not do.

Finding and Trusting the Data That Exists

Before any model can be trained or any hypothesis tested, researchers must identify where relevant data lives, whether they can access it, and whether it is trustworthy enough to use. AI can assist with retrieval and matching across heterogeneous sources, but it cannot substitute for a researcher's critical judgement about provenance, quality, bias, and ethical constraints. The speakers below both built their research around unusual data-access challenges — one harvesting a global stream of live radio, the other piecing together fragmented agricultural market signals — and both found that understanding the limits of their data was as important as acquiring it.



David Sears

Music cognition & computation — Mirage global radio database

Sears's project began with a single evening spent spinning a globe on Radio Garden, an API-connected streaming service covering nearly 40,000 live stations worldwide. The research question — whether it was possible to study musical diversity at a genuinely global scale — immediately became a data-access problem. He monitored 10,000 stations over three months, pulling metadata for a million streaming events, then matched each against open-access databases including Wikidata, MusicBrainz, and Spotify. Critically, he did not simply trust the matches: every field was assigned a reliability score, lower-confidence records were flagged rather than discarded, and the dataset was published in reliability quartiles on Zenodo so other researchers could choose their own quality threshold. His most important finding about the data itself — that open-access databases are systematically biased toward the global north, leaving African and Asian artists unmatched — became a research contribution in its own right.



Vijay Giri

UMich Dearborn — specialty-crop harvest decision support

Giri's work on forecasting systems for Michigan asparagus farmers exposed one of the starkest data-access failures in agricultural research: the official USDA data page for the crop had not been updated since 2022, import tracking had ended in 2021, and cost-of-production figures were nearly a decade old. Farmers making hourly decisions were relying on data that was years out of date. The project's most time-consuming phase was assembling a unified pipeline from heterogeneous public sources — USDA reports, cooperative pricing data, import tracking, satellite signals — each with different formats, different update frequencies, and no shared schema. "Getting them all into one unified pipeline and synchronising them independently was the most challenging part," Giri said. That infrastructure challenge, not the modelling, was where the project lived or died.

Cleaning, Structuring, and Augmenting Before Analysis

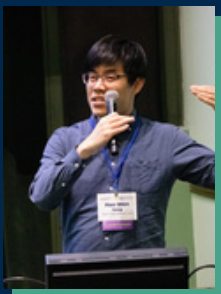
Raw data is almost never ready for a model. Researchers must clean inconsistencies, handle missingness, annotate examples, and make principled decisions about what counts as valid input — all while keeping the researcher in the loop rather than delegating these choices silently to automated pipelines. AI can assist with annotation and augmentation, but the transparency and validity of those choices shape everything that follows. The researchers below each faced severe data-preparation challenges: one confronting missing data across 40 hospitals, another dealing with a dataset so small it required creative synthetic augmentation.



Ali Namvar

Galban Lab — STREAM ICU monitoring framework

Namvar's team spent roughly 80% of the STREAM project's total effort on data preparation alone — a figure he cited not as a complaint but as a lesson. The two public ICU datasets used, EICU (158,000 patients) and MIMIC (19,000 patients), were rich but messy: data sparsity was common, variables were inconsistently defined across hospitals, and missingness in high-dimensional space created compounding challenges for a method like optimal transport, which is computationally expensive to begin with. Every decision about how to handle a missing value or reconcile an inconsistent label was a researcher judgement call with downstream consequences for which physiological states the algorithm would detect and how mortality risk would be attributed. Data quality, he concluded, is everything: "If you skip the pre-processing, no model is saving you."



Hao-Wen Dong

Performing Arts Technology — AI for a cappella singing practice

Dong's team faced a data-preparation problem with an unusual solution. Building a source-separation model for a cappella music — to isolate soprano, alto, tenor, bass, and beatbox tracks from a mixed recording — requires labelled multi-stem audio data that simply does not exist at scale. Their entire studio-quality dataset was 2.6 hours, far below what separation models typically need. After fine-tuning an existing music source-separation model on this small corpus with modest results, the team turned to AI-powered singing voice cloning to generate synthetic augmented training examples. The cloned voices were not perfect — the originals carry more emotion — but they were statistically diverse enough to meaningfully improve the model's generalisability. Preparation and augmentation, not architecture, was the bottleneck.

Choosing the Right Infrastructure, Not the Biggest One

Researchers tend to reach for the largest, most capable model available — but scalability, cost, data sensitivity, privacy, and reproducibility all demand that compute choices be made thoughtfully rather than by default. Cloud platforms offer scale; local infrastructure offers control and confidentiality. The decision between them is not merely technical: it shapes what research is ethically feasible, what can be audited, and what can be sustained over time. No speaker in the series engaged with this tradeoff more directly or rigorously than Zia Chi.



Zia Qi

School of Social Work — local AI for child welfare case narratives

Qi's research involves over 1.3 million case narrative records from Michigan's child welfare system — detailed, sensitive documents written by CPS workers describing substance abuse, domestic violence, housing instability, and child maltreatment. Sending that data to a commercial cloud API was never an option. Her Phase 4 journey was therefore a systematic benchmarking exercise: she tested models across a range of sizes on her actual classification tasks with her actual data, holding herself to explicit quality thresholds set before running the tests. The finding that surprised her was that models as small as 4 billion parameters could achieve near-perfect agreement with human coders on her tasks — making large cloud models unnecessary and local inference on a single workstation GPU entirely sufficient. Small is not a compromise, she argued; it is often the right answer. "Small and local go hand in hand. They are what makes it possible to do serious AI work with sensitive data on hardware you own."

Evaluating Early Results Before Committing to More

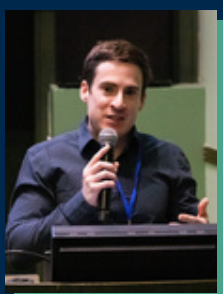
A baseline analysis is not a preliminary formality — it is the moment at which researchers discover whether their data and approach are pointing in a useful direction, and what the limits of simple methods reveal about the problem. Treating early outputs as definitive findings is one of the most common errors in AI-assisted research. The researchers below used baseline methods to expose structural problems — spurious clusters, cold-start biases, data-embedded artefacts — that more sophisticated approaches would only have amplified if left unaddressed.



Peter Bahr

Strada Institute — UMAP-based student typology for community colleges

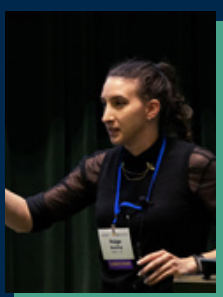
Bahr's journey begins with an uncomfortable baseline finding: K-means cluster analysis, the default method used across decades of research to categorise community college students, had been generating results that were largely spurious. The algorithm produces clean-looking groups whether or not genuine structure exists in the data, and in multi-dimensional behavioural data — where human beings cannot visualise distances directly — no one had been able to see through the illusion. UMAP (Uniform Manifold Approximation and Projection) entered the workflow not as a replacement clustering method but as a baseline diagnostic: a dimensionality-reduction tool that revealed the actual shape of the data before any clustering was applied. The finding was striking: roughly 70% of community college students, in terms of their behavioural patterns, cannot be cleanly segmented from one another. The baseline revealed that the decades of typological research had been building on sand.



Bryan Goldsmith

Chemical Engineering — ML for catalysis and materials discovery

Goldsmith's baseline methodology is Bayesian optimisation: train a Gaussian process model on an initial dataset of quantum mechanical calculations, estimate uncertainty across unexplored materials, and use an expected improvement acquisition function to select which DFT calculations to run next. The loop — calculate, train, select, calculate — is computationally modest by modern standards, but it was the baseline approach that mapped out which doped iridium oxide alloys were thermodynamically stable in acid before any synthesis work began. That baseline identified tantalum, tungsten, and molybdenum as promising dopants and flagged iridium-aluminium oxide as understudied — a finding later confirmed experimentally, with the material proving more active and more stable than pure iridium. The baseline was not a stepping stone to something fancier; it was the scientific result.



Paige Bowling

Computational Chemistry — ML to accelerate drug discovery via lambda dynamics

Bowling's first model — a contextual bandit trained to predict bias coefficients for multi-site lambda dynamics simulations — immediately revealed a classic baseline failure: the cold-start problem. A contextual bandit with no prior knowledge has no idea which direction to explore, and the initial runs produced results that were neither physically meaningful nor practically useful. The fix, behaviour cloning (pre-training the model on representative examples before live learning), is itself a baseline technique — a structured warm-up that gave the system enough orientation to begin improving. The subsequent discovery that her graph-based embedding scheme could not distinguish between ortho, meta, and para substituents — chemically distinct environments that the model was treating as identical — came only because the baseline runs were examined closely enough to notice the failure mode. "Problems do arise," she noted. "Some are inherent in the method you choose, and some you finally learn after you run some ML models that are embedded in your dataset." Those baseline failures drove every subsequent design decision.

Selecting and Evaluating Sophisticated Methods with Care

Advanced AI and machine learning methods offer genuine power, but that power comes with risks: hidden assumptions, poorly characterised uncertainty, sensitivity to distributional shift, and outputs that can appear authoritative while being subtly wrong. Researchers at this phase must not only select the right method for their objectives, but situate its outputs within domain knowledge and subject them to rigorous scrutiny. The speakers below span neurosurgery, medical imaging, computational chemistry, power systems engineering, and agricultural forecasting — each illustrating a different dimension of what responsible advanced analysis looks like in practice.



Todd Hollon

Neurosurgery — AI-guided brain tumour resection (*FastGlioma*)

Hollon's FastGlioma system uses a large encoder pre-trained on histological images, distilled into a slide-score model that estimates tumour infiltration at a surgical margin in real time. The advanced analysis challenge was not purely technical: the system's outputs had to be validated against the current clinical standard of care — intraoperative MRI and fluorescence-guided surgery — in a multi-site study across UCSF, NYU, and Vienna. The ROC curve comparison showed FastGlioma outperforming both existing options, and the rate of optimal resections increased while surgical errors decreased. But Hollon was careful about what this means: the device is not yet FDA-approved, the trial was non-interventional, and the sociological question of how surgeons change their behaviour when working alongside an AI system remains open and important.



Mathias Wilms

Radiology — Parkinson's disease classification from brain MRI

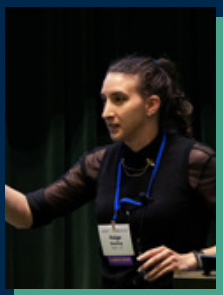
What appeared to be a straightforward binary classification problem — Parkinson's yes or no from a T1-weighted MRI — turned into a masterclass in algorithmic bias. A standard convolutional neural network trained on data from 40 centres worldwide achieved reasonable accuracy overall, but failed on held-out sites in ways that clinical saliency maps had not predicted. Wilms's team used a frozen-feature probing approach: if the features learned for Parkinson's classification also predicted scanner manufacturer, field strength, and acquisition site with equal or higher accuracy, the model had learned scanner artefacts rather than disease. It had. The response was two new research directions: a causal generative harmonisation model that removes site effects while preserving biological variation, and the Simba framework for synthetic data generation to systematically benchmark how and where biases propagate through network layers.



Rabab Haider

Engineering — AI optimisation proxies for power grid reconfiguration

Haider's group is trying to replace or accelerate a mixed-integer optimisation solver with neural network proxies for real-time power grid topology decisions — a problem where a physically infeasible output is not merely inaccurate but potentially dangerous. Her advanced analysis journey has been a systematic evaluation of different physics-embedding strategies: exploiting operational constraints from industry partners, building repair layers to project infeasible solutions back into the feasible space, and using graph neural networks with type-aware edges to encode the grid's physical topology directly into training. Each approach was tested across test systems ranging from 30 to 82,000 buses, and the benchmark results — including solving a case that Gurobi, the leading commercial solver, could not — demonstrate meaningful progress toward grid-scale deployment.



Paige Bowling

Computational Chemistry — ML to accelerate drug discovery via lambda dynamics

Bowling's journey through method selection is one of the most layered in the series. Starting with a contextual bandit, she navigated behaviour cloning to solve the cold-start problem, then added curriculum learning to progressively increase molecular complexity, then rebuilt her embedding scheme entirely after discovering it could not distinguish ortho, meta, and para substituents — three chemically distinct configurations that behave differently in energy space. The final architecture, using a spatial graph representation drawn from machine-learned molecular force fields rather than a collapsed 2D string, captured the physical geometry the earlier scheme had flattened away. Each pivot was grounded in chemistry, not trial and error: "I can't assume that traditional data pipelines will work perfectly for AI. I want to try the simplest representation first, but I have to keep the physics as part of the problem." The iterative evaluation of model assumptions — what each architecture can and cannot represent — is what makes her journey a textbook example of responsible advanced analysis in a domain where physical fidelity is non-negotiable.

Building Workflows That Can Be Reused, Sustained, and Scaled

Sometimes the research question cannot be answered with existing tools — the researcher must build the infrastructure themselves. Custom engineering at this phase means designing modular, reproducible AI-enabled workflows that integrate data, models, and computational resources in ways that others can build on. The speakers below represent three distinct flavours of this challenge: an automated laboratory that runs a million biological experiments, a cloud-hosted transcription API serving 40 institutions, and a citation-analysis pipeline processing 30 million paper pairs.



Pictured:
Benjamin David

Paul Jensen & Benjamin David

Biomedical Engineering — BactEriAI robot scientist

The Jensen lab's BactEriAI system applies reinforcement-learning principles — the agent designs experiments, receives data, updates an internal model, and plans the next cycle — to the question of what nutrients different bacteria need to grow. The engineering challenge, David explained, was not the model architecture but the wet-lab automation infrastructure required to close the loop: custom liquid-handling robot protocols written by the Perfecto software, plate-scheduling logic, robotically managed incubation chambers, and a data pipeline capable of returning growth profiles to the agent in a form it could learn from. Starting at 300 experiments per day in 2020, the lab scaled to 10,000 per day within a year. The system has now completed over a million experiments and is being extended — in partnership with the Alian Foundation — to culture a thousand microbial species in a thousand conditions each, building what would be the world's largest open dataset for microbial phenotyping.



William Weaver

MIDAS — Voucher Vision herbarium digitisation platform

Natural history collections hold roughly 400 million dried plant specimens globally — a vast store of ecological and evolutionary knowledge that has been accumulating for 400 years. The bottleneck is transcription: getting the handwritten label data off the paper and into searchable databases. Weaver's Voucher Vision platform combines optical character recognition with vision-language models and structured prompt templates, deployed as a cloud-hosted API on Google Cloud. The workflow is explicitly modular: institutions can select which fields to extract, which model to use, and how aggressively to review AI-generated outputs. With 40 collaborating institutions and a quarter of a million specimens transcribed in six months, the new bottleneck is no longer transcription speed but the databases' capacity to handle AI-assisted data and the editing interfaces needed to maintain quality. The platform also builds in digital repatriation — bidirectional translation of label text back into the language of the country where the specimen was originally collected.



Hong Chen

School of Information — citation fidelity measurement at scale

Chen's research on how accurately scientific claims are conveyed as they move through citation chains required a custom three-step NLP pipeline: a fine-tuned BERT model to identify reporting citations (distinguishing them from background or methodological ones), a second fine-tuned model to extract the corresponding results and conclusions from the cited paper, and a third specialised metric to score how much information changed between the original claim and its reported version. Applied to 30 million citation pairs, the pipeline produced the first large-scale empirical evidence for the "telephone game" effect in science — showing that fidelity declines with temporal distance, that intermediary citations compound distortion when they are themselves unfaithful, and that self-citations are the most faithful of all. The pipeline was designed to run on standard cluster hardware, making the findings reproducible without requiring commercial LLM APIs.

Validating, Auditing, and Making Results Trustworthy

An AI-generated output is not a finding until it has been validated against theory, existing knowledge, or empirical reality — and documented in a way that others can scrutinise and reproduce. This phase is where research earns its credibility. The speakers below each subjected their AI outputs to rigorous external validation: multi-site clinical trials, prospective experiment on held-out datasets, controlled user studies, and systematic human-judge agreement testing.



Todd Hollon

Neurosurgery — FastGlioma multi-site validation

The reproducibility test for FastGlioma was built into the project from the start: the slide-score model was deployed and evaluated independently at three international institutions — UCSF, NYU, and Vienna — without retraining at each site. Consistent performance across institutions, tumour types, and patient populations (adult and paediatric) was a prerequisite for any clinical credibility. The ROC curve published in the New York Times-noted paper compared FastGlioma not against an arbitrary baseline but against the two methods patients would actually receive at a university hospital. The finding that optimal resection rates increase and surgical errors decrease is meaningful precisely because it is framed in clinical outcomes, not model accuracy metrics alone.



Farnaz Jahanbakhsh

CSE — personalised content moderation system

After building the browser extension that provides real-time personalised transformation of social media content, Jahanbakhsh ran a controlled experiment to verify that the transformations the AI pipeline selected actually matched what users would have chosen for themselves — closing the loop between system design and user preference. The qualitative feedback from users was striking in its own right: participants reported that far from encouraging avoidance of difficult content, the tool enabled re-engagement — they were able to stay connected with people and topics they had previously been forced to mute entirely. Outputs were audited not just for technical correctness but for real-world psychological effect, and the results were validated against the formative study's design goals



Jingyi Qiu

School of Information — measuring AI hype in scientific papers

Qiu's persona-based rhetorical scoring framework required two layers of validation before it could be trusted as a measurement instrument. First, the personas were tested to confirm they genuinely spanned the full rhetorical distribution — by having an LLM judge compare persona-generated abstracts against real ones, and checking that win rates ran from near-certain to near-zero across the 30 personas. Second, the judges' ratings were compared against human expert assessments to confirm inter-rater reliability. Only after both validations did the team apply the pipeline at scale to identify the sharp rise in rhetorical inflation since 2023 and gather evidence linking it to LLM-assisted writing. The tool that was built to measure AI hype had itself to be audited for the hype it might introduce.

Sharing Data, Models, and Insights Across Communities

Research that stays inside a lab contributes only a fraction of its potential value. Phase 9 is about building the infrastructure, norms, and communities through which AI-enabled findings, datasets, models, and workflows reach the people who can use and build on them. The speakers below represent two very different scales of knowledge sharing — one building a global data infrastructure for music research, the other co-founding the communities of practice that have brought African researchers into the centre of global AI development.



David Sears

Music cognition — Mirage Project open dashboard and dataset

The Mirage Project was designed from the outset as shared infrastructure, not a private dataset. The million-event corpus of radio metadata is open access, published on Zenodo in reliability-stratified layers. The online dashboard at MirageProject.org allows researchers without a developer background to explore, filter, and export data directly. A Python client library is nearing completion for programmatic access. And GeoListen — a geographic music-guessing game in the iOS and Google Play stores — transforms the general public into research participants, collecting data on how listeners across the world form geographic associations with unfamiliar music. Sears's explicit goal is to create infrastructure that other researchers can build on without having to repeat the 18 months of data collection his team invested.



Farnaz Jahanbakhsh (Deep Learning Indaba & Masakhane)

University of Pretoria / co-founder, Deep Learning Indaba and Masakhane

Jahanbakhsh's introductory remarks described the grassroots knowledge-sharing infrastructure she and others built to address a structural gap: African researchers were largely absent from global AI development, and African languages were almost entirely absent from NLP research. Deep Learning Indaba — now the leading AI research gathering on the African continent, with 1,200+ registered participants from 50+ countries — began in 2016 as a meeting of 300 people in Johannesburg. Masakhane, a spin-off focused on NLP for African languages, has grown to 3,000 researchers and won four best paper awards at ACL 2025. The Masakhane African Languages Hub now runs its own grant-making programme. These are not just community events — they are the knowledge-sharing infrastructure without which no amount of model development would reach the researchers and languages that need it most.

CONCLUSION

A LIFECYCLE, NOT A CHECKLIST

The nine phases of the AI Research Lifecycle are not a recipe to follow in order. Real research moves through them unevenly — looping back from Phase 6 to Phase 1 when a model's outputs reveal a question that was never properly asked, jumping from Phase 2 straight to Phase 7 when the data infrastructure has to be built before any data can be found. What the lifecycle offers is not a roadmap but a vocabulary: a set of named moments at which researchers need to stop, think, and make deliberate choices about what AI is doing in their work and what it is not.

The researchers in this booklet span a remarkable range of disciplines, methods, and problems. But a few themes recur across almost every talk. Domain expertise is never optional — the most technically sophisticated systems in this collection work because they were built by people who understood the underlying problem deeply enough to know when the AI was wrong. Multidisciplinary collaboration is not a nice-to-have; it is the mechanism by which AI outputs get validated against reality. And honest accounts of failure — the contextual bandit that could not handle a cold start, the K-means clusters that turned out to be spurious, the Parkinson's classifier that learned scanner artefacts instead of disease — are as scientifically valuable as the successes.

AI is changing every phase of the research lifecycle. The researchers at Michigan are not waiting to see how that change lands — they are making it, carefully and in public, one journey at a time.



**MICHIGAN INSTITUTE
FOR DATA & AI IN SOCIETY**
UNIVERSITY OF MICHIGAN

Michigan Institute for Data & AI in Society
Weiser Hall, Suite 600
500 Church Street Ann Arbor, MI 48109

midas.umich.edu · midas-contact@umich.edu