

Opportunities for ML/AI to Derive Better Outcomes in Hope

**Presentation to the Coventry University
Health and Community Wellbeing Theme.**

Prof. Matthew England and Dr Michael Ajao-olarinoye



Who are we?



From Coventry University Research Centre for Computational Science & Mathematical Modelling

Prof. Matthew England:

Professor of Computer Science; Director of CSMM
(specialist on integration of AI with mathematical software).



Dr Michael Ajao-olarinoye:

Research Fellow; PhD in Computer Science (projects in pandemic modelling; aneurysm modelling).

Getting to know HOPE



We met Gabriela and Andy at one of their presentations about HOPE at Coventry University in early 2025. We discussed the potential for ML/AI to enhance the HOPE Program.

We collaborated soon after to apply for the **HOPE-MOVE project** (2025-26)

The project as a whole, is expanding the HOPE Program to those living with osteoarthritis and chronic pain in the hip or knee.

We lead a work package trialing ML/AI tools on Hope data.

First we acknowledge our collaborators.



Gabriela Matouskova

CEO AND NON-EXECUTIVE DIRECTOR

Gabriela leads on business development and growth. She loves all things crossing arts, health, creativity, innovation, community and research. Gabriela cared for her mum, who was diagnosed with terminal cancer in 2013. Her personal experience of caring and loss fed into development of our Hope Programme for Carers.



Prof Andy Turner

NON-EXECUTIVE DIRECTOR, RESEARCH

Andy is the founding member of our social enterprise, Professor of Health Psychology at Coventry University, and co-inventor of the Hope Programme. He has developed, delivered and researched self-management programmes for over 25 years. Andy lives with grade 4 arthritis in both knees and manages his pain through playing tennis.



Dr Bethan Alper

QUALITY ASSURANCE LEAD

Bethan leads on training, online facilitator assessments and quality assurance for the Hope Programme. She loves music, dancing and fancy dress. After completing her PhD at Coventry University, Bethan worked in healthcare until her Multiple Sclerosis diagnosis in 2018. She completed the Hope Programme in 2020 and now leads our quality assurance and facilitator training.



Elizabeth Horton. Deputy Director (Quality),
CU Health and Care. Previously Associate
Director for Research and Engagement



Matthew Jones. Managing Director
at Magic Square Systems.

Introduction to HOPE

(Thanks to Andy Turner and his AI)

TECH FOR GOOD: TECHNOLOGY SOCIAL ENTERPRISE OF THE YEAR 2024



36,000

HOPE PROGRAMME
PARTICIPANTS TO DATE



16

EVIDENCE BASED
CO-PRODUCED
SELF-MANAGEMENT DIGITAL
INTERVENTIONS



5

STAFF
80% FEMALE, 60% LIVING WITH
LONG TERM CONDITIONS



15

PEER-REVIEWED ACADEMIC
JOURNAL ARTICLES PUBLISHED

powered by



£558,000

CONFIRMED VALUATION OF
OUR DIGITAL PLATFORM



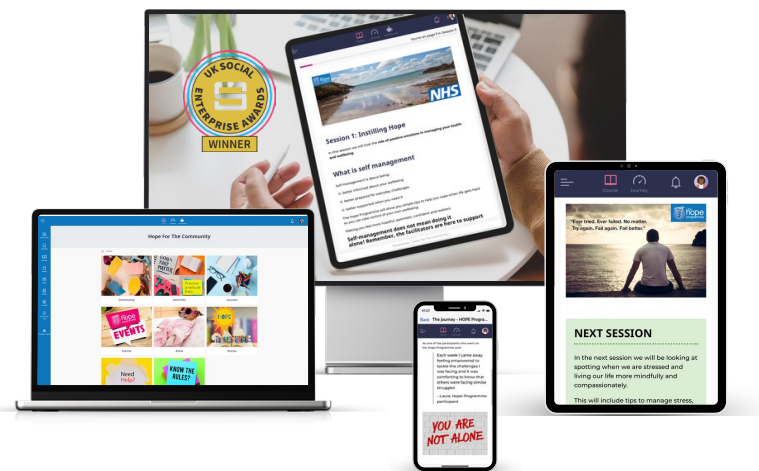
£81K

PROFIT RE-INVESTED
INTO OUR COMMUNITIES



The Hope Programme empowers people living with long-term conditions and carers to take control of their health and wellbeing closer to home.

- ✓ **EVIDENCE BASED** courses built on 25+ years of research
- ✓ **MULTI CONDITION** holistic support focused on the person
- ✓ **CO-DESIGNED** and **PROVEN** to improve mental wellbeing, patient activation and quality of life
- ✓ **SCALABLE** digital-first with local in-person delivery



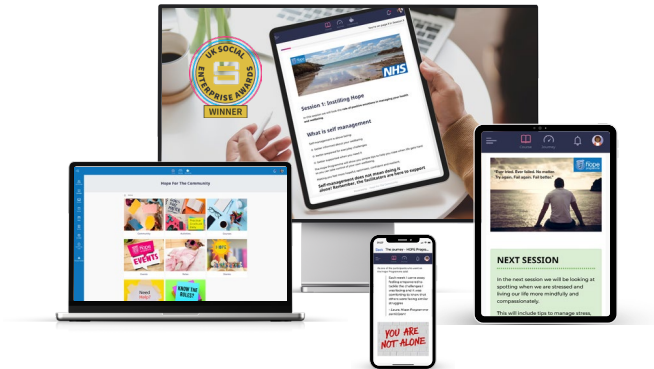
36,000

PEOPLE SUPPORTED TO DATE

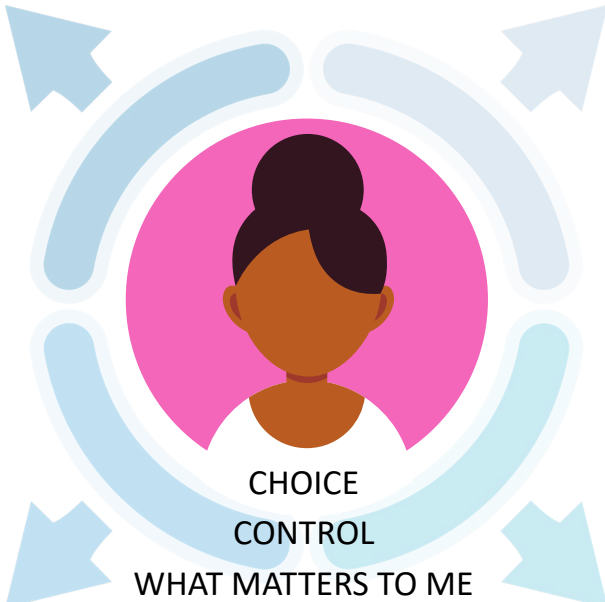
EMPOWER PEOPLE: PERSON-LED + HOLISTIC + FLEXIBLE COURSES



DIGITAL
at home
anonymous
asynchronous
groups or self-guided
online community



adult (18+)
any long-term condition(s)
unpaid carer
low level mental health needs
lonely and isolated
low motivation and/or confidence



IN-PERSON
in community
groups
local facilitators
trusted partners

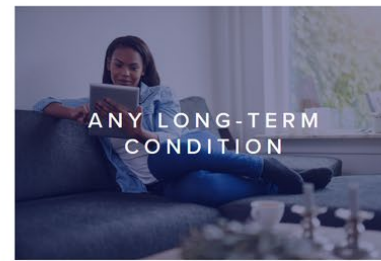


HOPE PROGRAMME COURSES OFFER

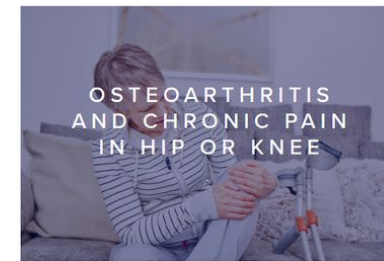
GROUP FACILITATED COURSES online
or in-person - start on set dates -
meet others in a similar situation -
get support from trained facilitators

SELF-GUIDED COURSES online - start
now - learn on your own - technical
support

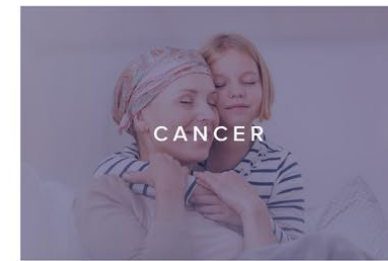
ONLINE COMMUNITY ongoing course
content access and self-management
- peer-support - events and more



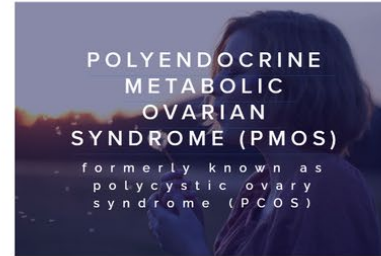
ANY LONG-TERM
CONDITION



OSTEOARTHRITIS
AND CHRONIC PAIN
IN HIP OR KNEE



CANCER



POLYENDOCRINE
METABOLIC
OVARIAN
SYNDROME (PMOS)
formerly known as
polycystic ovary
syndrome (PCOS)



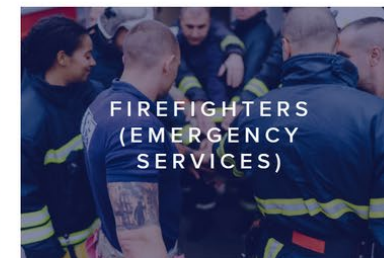
MULTIPLE
SCLEROSIS



LONG COVID



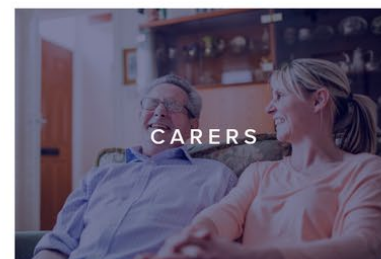
IDIOPATHIC
INTRACRANIAL
HYPERTENSION



FIREFIGHTERS
(EMERGENCY
SERVICES)



WORKPLACE
WELLBEING



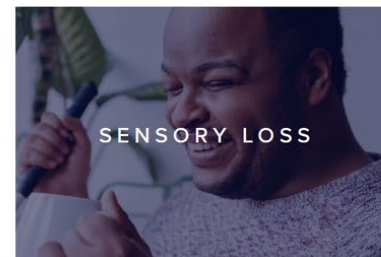
CARERS



PARENTS OF
AUTISTIC
CHILDREN



PARENTS OF
CHILDREN WITH
CANCER



SENSORY LOSS



LEARNING
DISABILITY



DEMENTIA

Redefining self-management through strengths, not deficits

	Traditional Medical Model	The HOPE Model
Core Focus	Pathology, symptom management, and breaking the negative circle of fear and anxiety. 	Positive psychology, gratitude, and building the Upward Spiral of resilience (Broaden-and-Build theory). 
Psychological Mechanism	Self-Efficacy: Focused on specific, isolated behaviours and agency beliefs. 	Hope Theory: Focused on enduring, cross-situational goals combining both motivation and planning. 
Group Dynamic	Expert-led, didactic instruction. 	Peer-facilitated group curative factors (Universality, Altruism, Instillation of Hope). 

SNYDER ADULT STATE HOPE SCALE (ASHS)

Hope is a cognitive set based on a reciprocally-derived sense of successful agency and pathways.



AGENCY SUBSCALE

Goal-Directed Determination (Willpower)

The motivational element focused on commencing and persevering with goal pursuits.

1. At the present time, I am energetically pursuing my goals.
2. Right now, I see myself as being pretty successful.
3. At this time, I am meeting the goals that I have set for myself.



PATHWAYS SUBSCALE

Planning to Meet Goals (Waypower)

The perceived ability to produce plausible routes and alternative strategies to reach goals.

1. There are lots of ways around any problem that I am facing now.
2. If I should find myself in a jam, I could think of many ways to get out of it.
3. I can think of many ways to reach my current goals.

SCORING & INTERPRETATION (Based on an 8-point Likert scale: 1 = Definitely False to 8 = Definitely True)

Agency Score: 3 – 24

Pathways Score: 3 – 24

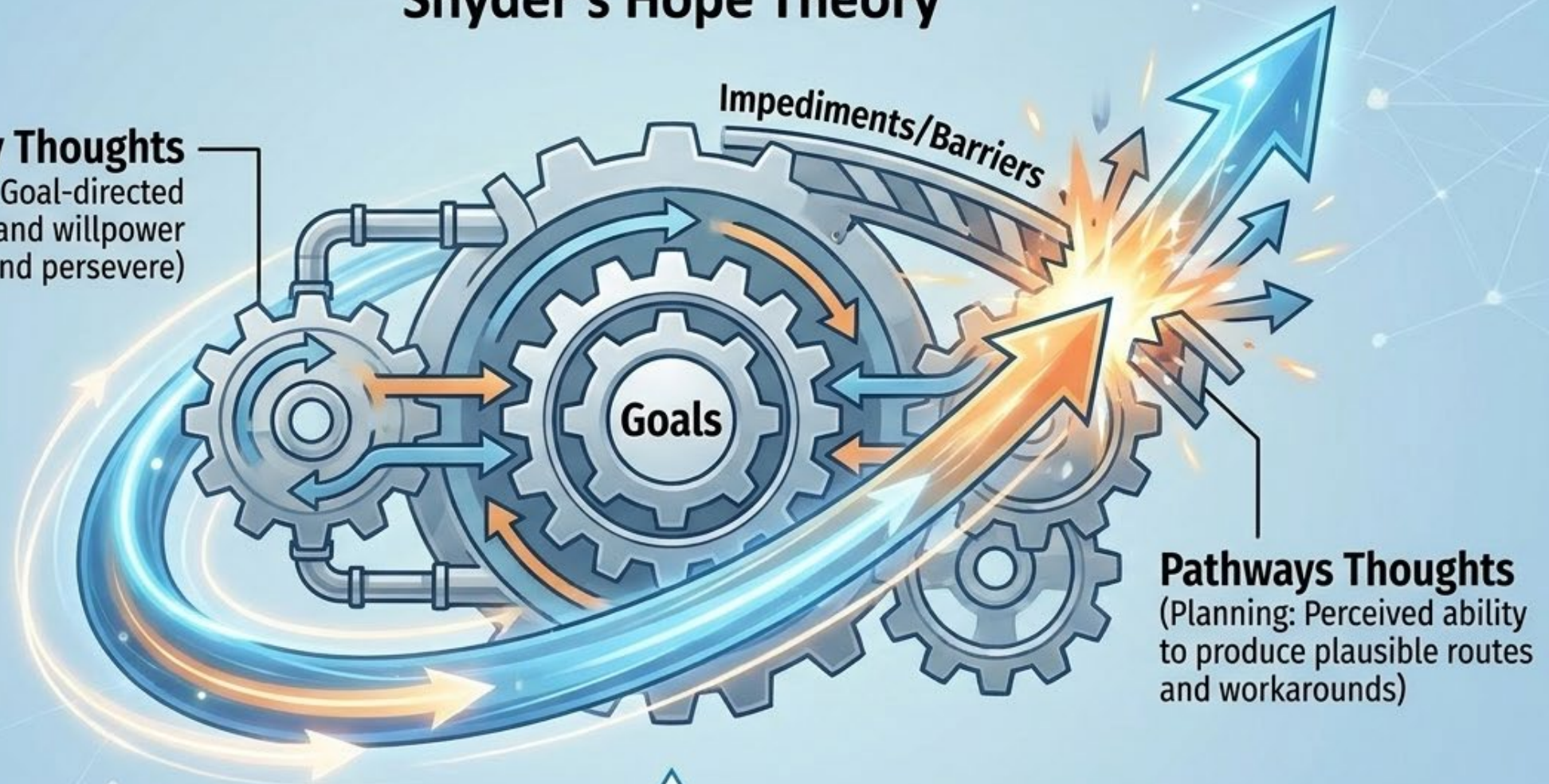
Total ASHS Score: 6 – 48

(Note: Higher scores indicate greater levels of hopeful thinking and a stronger capacity for self-management.)

The cognitive engine driving the intervention: Snyder's Hope Theory

Snyder's Hope Theory

Agency Thoughts
(Motivation: Goal-directed determination and willpower to commence and persevere)

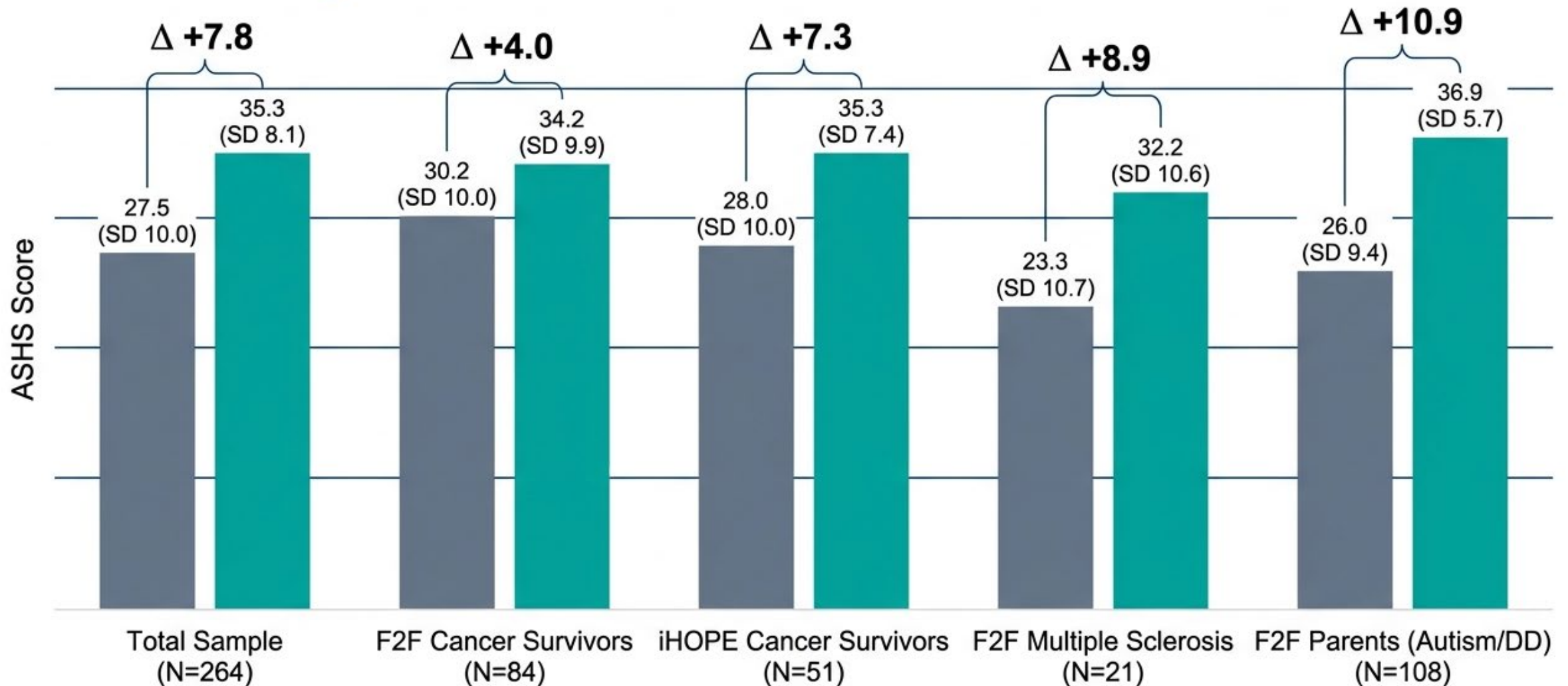


Pathways Thoughts
(Planning: Perceived ability to produce plausible routes and workarounds)

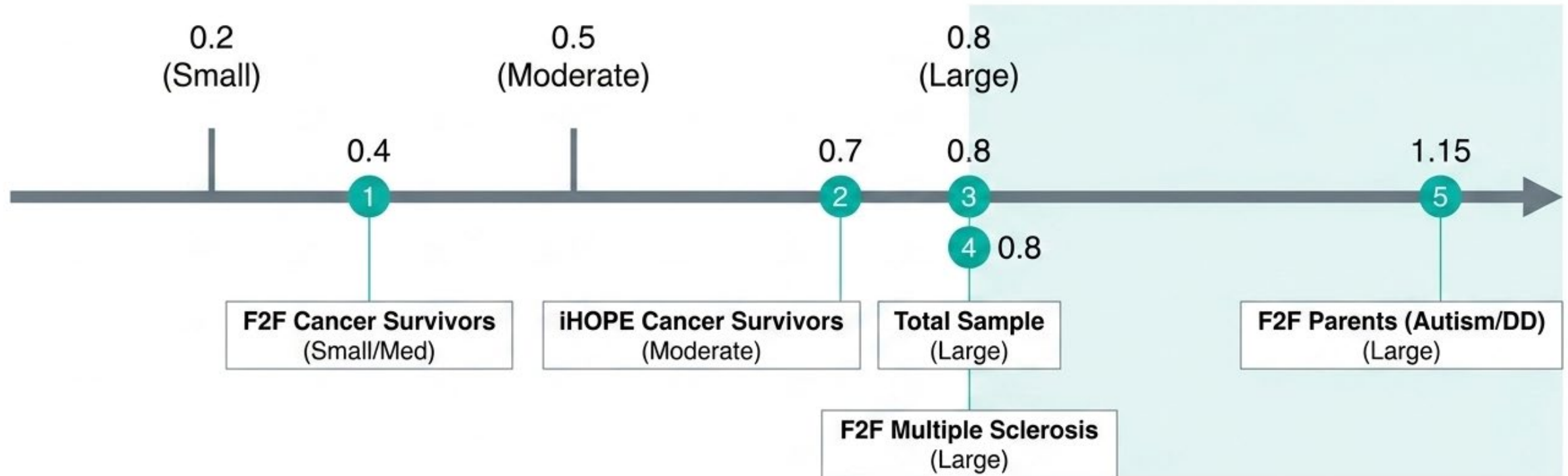


Over 264 participants combined across 4 feasibility studies demonstrated statistically significant improvements in the Adult State Hope Scale (ASHS), validating this mechanism.

Consistent Upward Trajectories in Adult State Hope Scores



Standardising the Impact: Magnitude of Change Across Cohorts



The Measurement Imperative

Whilst dominant models in clinical psychology focus on psychopathology and skills deficits, successful chronic care requires tracking the emergence of capability.

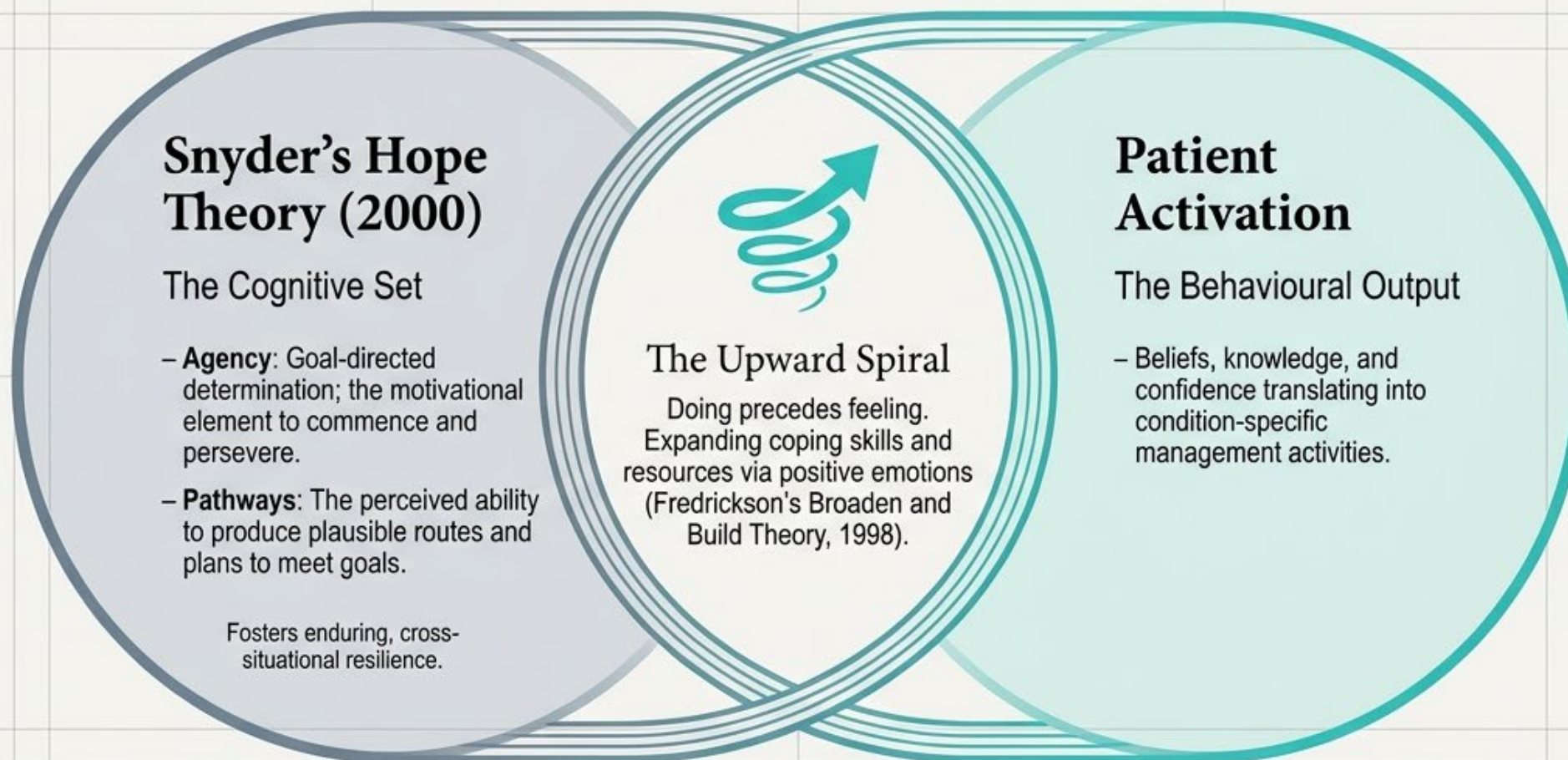
To scale self-management interventions like the HOPE Programme, we must empirically measure the transition from passive distress to active self-regulation.

The Catalyst for Self-Management

“Confidence about the ability to self-care is crucial to self-management: individuals with lower confidence in their self-care abilities are less likely to seek help, follow guidance, and perform self-management behaviours.”

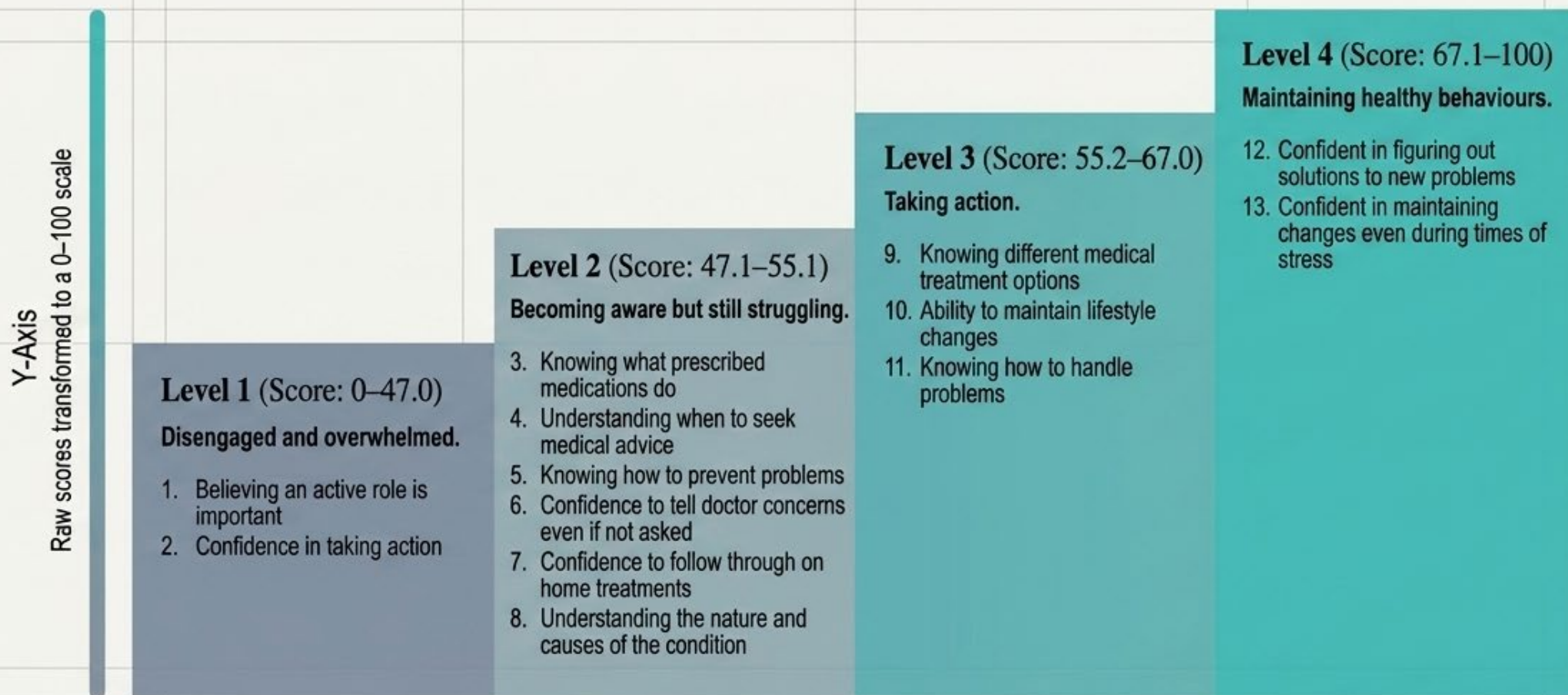
— Derived from Hibbard & Gilbert (2014)

The Cognitive Architecture of Activation



The Patient Activation Measure acts as the clinical barometer for the iterative, additive cognitive shifts fostered by hope-based interventions.

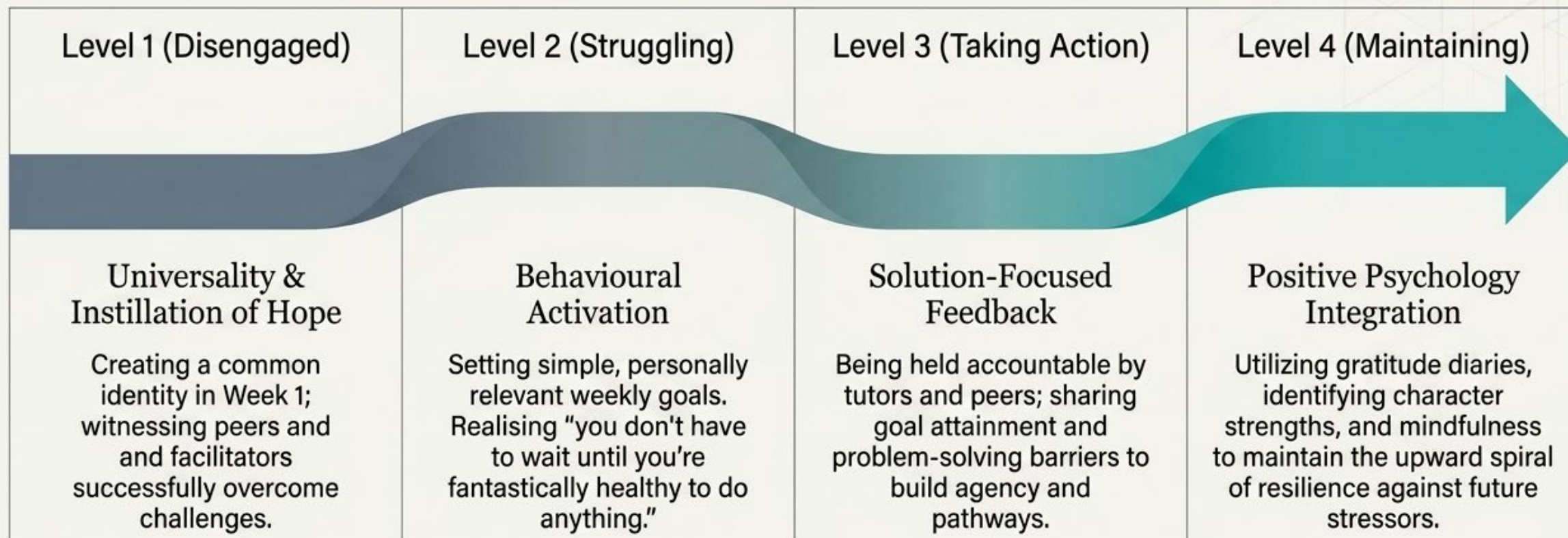
The Patient Activation Measure (PAM-13)



Note: The specific items and scoring thresholds presented are based on established clinical standards (Hibbard et al.) conceptually cited but not fully detailed in the primary source material.

Scaling Hope: Operationalising the Activation Gradient

Action Plan



The HOPE Programme leverages positive psychology not merely to alleviate distress, but to actively engineer the cognitive and behavioural leaps required to sustain Level 4 activation.

PATIENT ACTIVATION

Activation - the **knowledge**, **skills** and **confidence** in self-managing health leads to

- better health outcomes
- lower cost of care
- better patient experience



Phreesia

PATIENT ACTIVATION MEASURE® (PAM®) N=518*

*collected from Hope Programme courses delivered Sep 2024 - Apr 2026

PAM levels	Pre HOPE	%	Post HOPE	%
Level 1	211	41%	96	19%
Level 2	120	24%	109	20%
Level 3	155	29%	193	37%
Level 4	32	6%	120	24%

Before attending HOPE **64% people were at levels 1-2**, indicating that the people who would likely benefit the most from self-management were being recruited. **36% people were at levels 3-4**. Post HOPE there were **24% fewer people at levels 1-2 (40%)** and **there were 24% more people (60%) at levels 3-4** (taking and maintaining action to self-manage their condition). The average PAM score before HOPE was 51.8 changing to 60.4 post HOPE, an average increase of 8.6 points.

"I'm now having a more active role in my healthcare and looking after my mental health."

- Emma, 6 months after completing the Hope Programme



PATIENT ACTIVATION

EVALUATION - PRE AND POST COURSE (N=513)* EXPECTED IMPACT ON SAVINGS AND RETURN ON INVESTMENT

ALL LTC GROUP and SELF HOPE(exc outliers) Pre HOPE PAM level	People who completed pre and post HOPE (latest) PAM surveys	Change in PAM score per person (from pre to post)	Assumed cost saving per single- point PAM score change	Expected cost saving per patient	Total expected savings across all people (savings per person multiplied by number of patients)
Level 1	211	10.69	£317.00	£3,388.73	£715,022.03
Level 2	120	9.01	£174.00	£1,567.74	£188,128.80
Level 3	155	6.78	£171.00	£1,159.38	£179,703.90
Level 4	32	-0.09	£190.00	-£17.10	-£547.20
All people	518		Average pp	£2,089.40	£1,082,307.53

Annual
healthcare
costs of
someone living
with 2 LTCs is
£3,000 - Hope
shows
potential of
around **70% in
modelled
savings - costs
avoided** - per
person per
year.

*Phreesia Impact of the Patient Activation Measure (PAM®) July 2024 reference Self-management capability in patients with long-term conditions is associated with reduced healthcare utilisation across a whole health economy: cross-sectional analysis of electronic health records | BMJ Quality & Safety [Self-management capability in patients with long-term conditions is associated with reduced healthcare utilisation across a whole health economy: cross-sectional analysis of electronic health records | BMJ Quality & Safety](#)

**Matched pre/post PAM data, excluding outliers (surveys with too many 'not applicable' answers or uniform response patterns).

The Warwick-Edinburgh Mental Well-being Scale (WEMWBS)

Full 14-Item
WEMWBS

7-Item Short Version
(SWEMWBS)

1. I've been feeling optimistic about the future

2. I've been feeling useful

3. I've been feeling relaxed

4. I've been feeling interested in other people

5. I've had energy to spare

6. I've been dealing with problems well

7. I've been thinking clearly

8. I've been feeling good about myself

9. I've been feeling close to other people

10. I've been feeling confident

11. I've been able to make up my own mind about things

12. I've been feeling loved

13. I've been interested in new things

14. I've been feeling cheerful

The Baseline Reality: Profiling 942 Long COVID Patients

942

78.8%
Female

Total Participants

The Assessment:

Short-Form Warwick-Edinburgh Mental Wellbeing Scale (SWEMWBS)

The Metric:

Measures positive mental wellbeing on a scale of 7 to 35.

The Threshold:

A shift of just 1 point constitutes a "meaningful, minimal detectable change" (MDC).

1-Point
Minimal
Detectable
Change

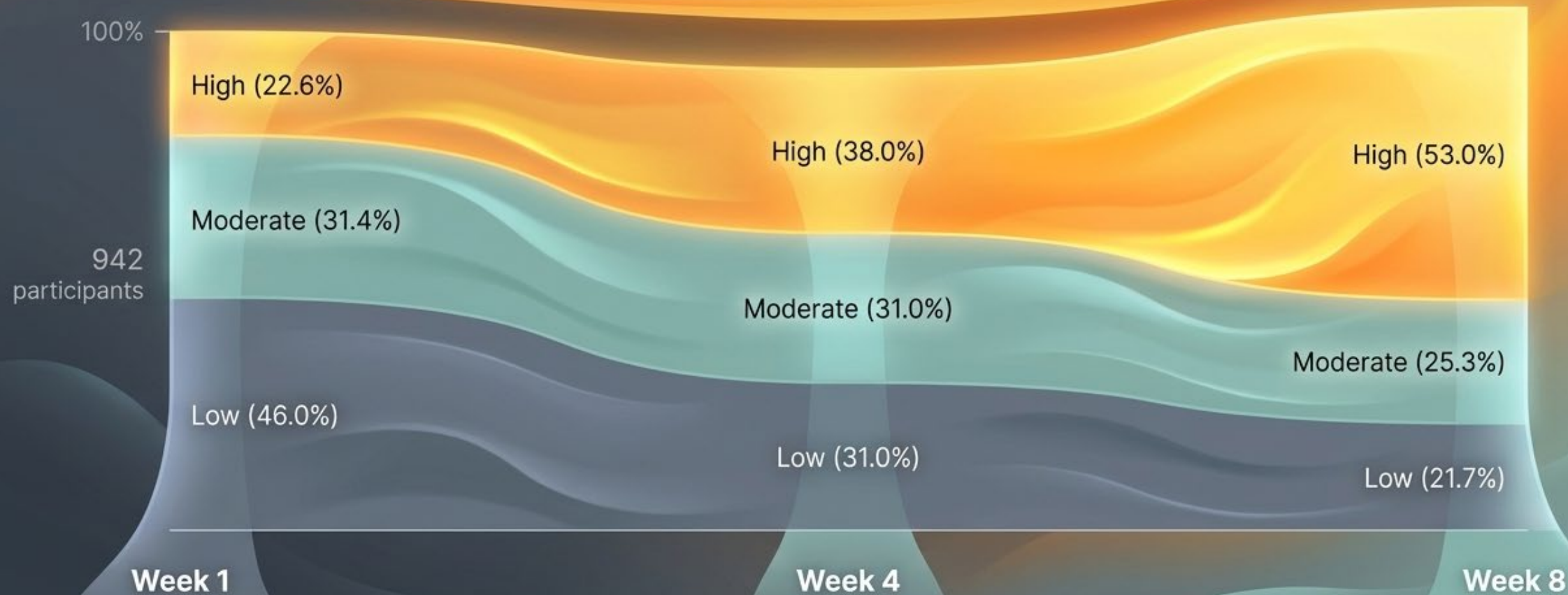


A Linear Trajectory of Meaningful Recovery



The Migration to High Wellbeing

The proportion of individuals experiencing 'High' mental wellbeing more than doubled.



Critical appraisal of the current evidence base



Strengths (Traction)

- Feasibility and acceptability proven across multiple delivery formats (Face-to-face, digital iHOPE).
- Strong, consistent signals of efficacy across a wide diversity of physical and psychological conditions.
- High usability scores (e.g., SUS score of 70.1; ORCHA certifications).



Limitations (Focus for Growth)

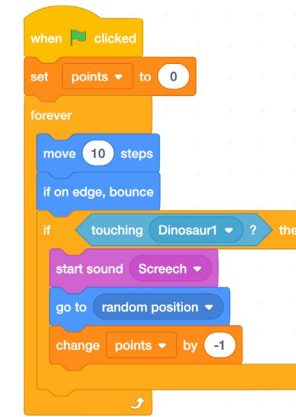
- Lack of active control groups (Waitlist controls cannot rule out expectancy/placebo effects).
- Small sample sizes in feasibility pilot trials.
- Highly homogenous, self-selected demographics (Predominantly White, female populations).

Crash Course on AI and ML

Crash Course: AI vs ML

Artificial Intelligence (AI) is the field of computer science that aims to create systems capable of performing tasks that typically require human intelligence.

- Includes e.g. rule-based systems; knowledge graphs.
- Recent success has come from Machine Learning.



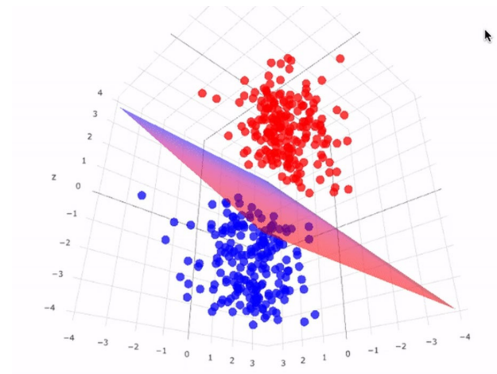
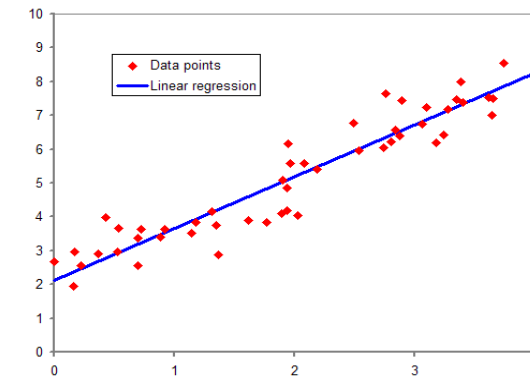
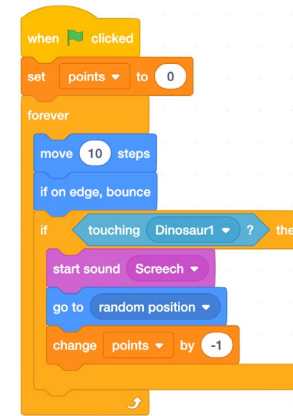
Crash Course: AI vs ML

Artificial Intelligence (AI) is the field of computer science that aims to create systems capable of performing tasks that typically require human intelligence.

- Includes e.g. rule-based systems; knowledge graphs.
- Recent success has come from Machine Learning.

Machine Learning (ML) focuses on systems to learn from data, instead of being explicitly programmed with rules.

- **Supervised learning** requires large quantities of data to learn a model from (provided by our world's increased digitalisation).
- **Unsupervised learning** requires a way to reward models based on outputs.



Crash Course: What about LLMs?

Large Language Models (LLMs) are often what people refer to by AI in 2026. These are an extreme form of machine learning:

- Models trained on all data available on the internet (at huge computational cost).
- Trained for the next word prediction task. E.g. given "Hope For the Community CIC is based in" what is the next word?
An LLM will likely predict "Coventry". Not because it has a file on H4C, but because this fact was present somewhere on the internet and so it is a likely next word.



	Traditional ML	Large Language Models (LLMs)
Core Objective	Predict & Classify new data based on patterns in training data.	Generative: To generate new objects based on patterns in training data
Input Data Type	Traditionally data is structured (e.g. spreadsheet) but can now operate on unstructured data.	Designed for unstructured data
Output Data Type	A label, score, class, or numeric value.	A new object, e.g. text, image, code.
Underlying technology	A wide variety to choose from which vary in power, computational requirements, data requirements and interpretability.	Deep neural networks (Transformer-based architectures). Huge, expensive to create, not interpretable.
Scope	Specialists: Model is usually trained for just one task.	Generalists: The same model can be applied on many different tasks.
Explainability	Low-High: Depending on the underlying technology.	Low: Although you can ask an LLM to explain its reasoning this is not guaranteed to be accurate!
Accuracy	High: If trained correctly, they yield consistent, reliable results	Variable: Susceptible to "hallucinations" (confident but false assertions) due to probabilistic text generation.
Computational Cost	Low to moderate: many models can be trained on a standard laptop.	Massive: Requires hundreds of GPUs, millions of dollars, and massive power consumption. Commonly obtained as a cloud service.

Initial Ideas for Working AI/ML into HOPE



HOPE already sit on a lot of data from previous users of the HOPE Program. We planned to see if some more value could be extracted from this data.

- Identify patterns of engagement from the data.
- Use these to predict current users at risk of dropping out.

These tasks can be address by traditional ML and Data Science.

The Hope Program is supported by facilitators who work to keep participants engaged. We also explored whether LLMs could shoulder some of that weight by suggesting communications for facilitators.



Bethan Alper



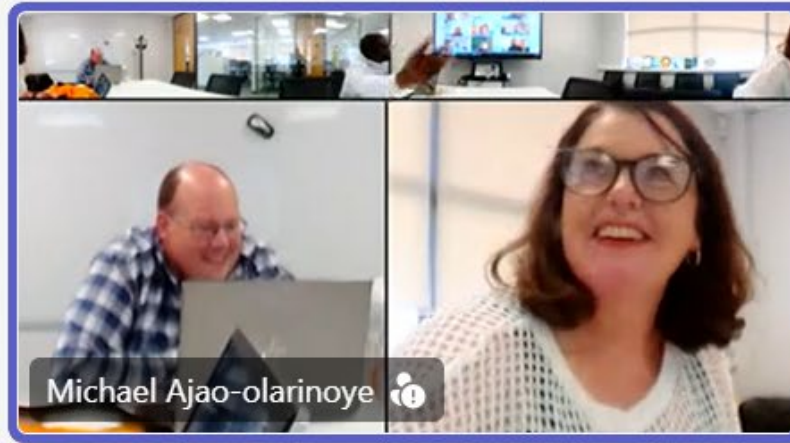
Suzz Keith



Jen Scoble



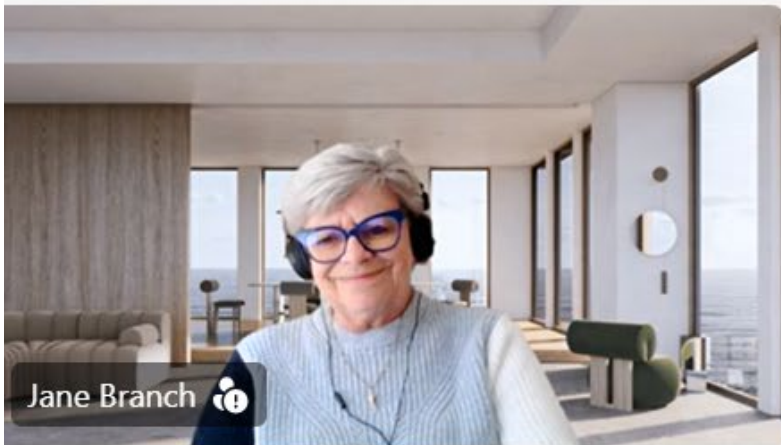
Elizabeth Horton



Michael Ajao-olarinoye



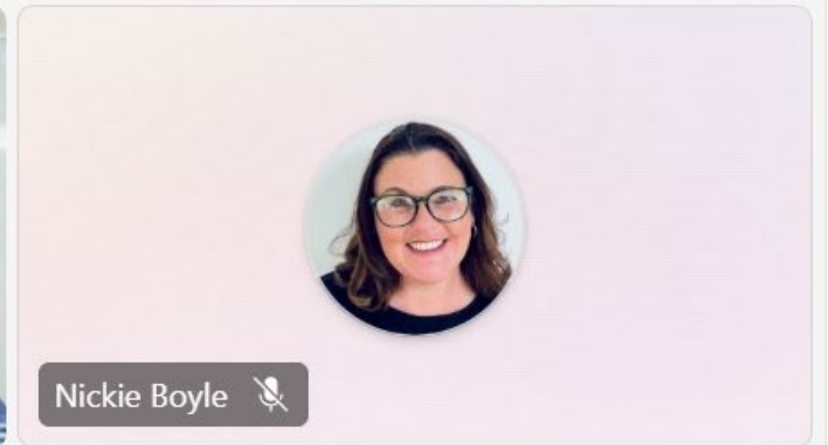
John Clare



Jane Branch



Matthew England



Nickie Boyle

Data Security and Research Ethics

This project has Research Ethics approval from Coventry University.

Key concerns were around security of personal data. Ensured this as follows:

- We were provided with anonymized data (participants were identified by a number rather than a name).
- Data kept locally on campus computers throughout.
- All ML models were local.
- Even out LLM experiments used only open-source LLMs running on large computers in Coventry; no data was shared with cloud models.
(This also had the advantage of being free to use!)

Proof of Concept Dashboard

SCORE AT WEEK W1 W2 W3 W4 W5 W6 (day 14) Σ Random Forest · T1.4 · AUC 0.94

Follow-up queue

62 <<

Search participants...

All Needs attention

Check in soon On track

P17	Needs attention	86%
Last active 14 days ago		
P19	Needs attention	86%
Last active 14 days ago		
P20	Needs attention	86%
Last active 14 days ago		
P21	Needs attention	86%
Last active 14 days ago		
P22	Needs attention	86%
Last active 14 days ago		
P24	Needs attention	86%
Last active 14 days ago		
P26	Needs attention	86%
Last active 14 days ago		
P31	Needs attention	86%
Last active 14 days ago		
P33	Needs attention	86%
Last active 14 days ago		
P35	Needs attention	86%
Last active 14 days ago		

Showing 1-10 of 62
← 1 2 3 4 5 6

P7

P7 Cohort 1600 · 17904

On track

< >

Snooze 7d

Dismiss



WHAT'S GOING WELL FOR P7

Browsing programme content	↓ 63%
Returning across multiple days	↓ 19%
Broad activity range	↓ 18%

↑ raises this person's risk · ↓ lowers it. Bar width = share of the explanation.



ACTIVATION LEVEL

High

Steady momentum — keep the touch light.



RECOMMENDED APPROACH

- Continue light-touch monitoring; no outreach required this week
- Acknowledge their progress at the next group touchpoint



WELLBEING SUPPORT CUE

Light-touch encouragement keeps momentum without overloading their inbox.

ENGAGEMENT SIGNALS

LAST ACTIVE

2 days ago

DISCUSSION POSTS

1

vs prior 14d: n/a

FACILITATOR CONTACT

4

to date

RECENT ACTIVITY

COLLAPSE

SHOW

Outreach · P7

Drafted by Qwen3 1.7B

PARTICIPANT POST

from 453d ago

GoalSetting

I am going to write a blog this week - topic TBC. on I will use my tablet . I will aim to do this Once let week

Regenerate

Write my own

SUGGESTED TONE

Warm check-in

Next-step nudge

Goal-focused

To: P7 · in-app message

Editable draft

P7, thank you so much for sharing your goal. We look forward to seeing all that you'll achieve in weeks ahead! 😊

113 chars

Helpful? 👍 🗑️ ⋮ Send

What this draft is based on

Proof of Concept Dashboard Produced

Hope Facilitator Assistant

localhost:3000/cohorts/1600

Summarize

Model

Qwen3.5 4B

47 need follow-up

40 high priority

0 contacted this session

hope-move

Participant support

IIH-COH10-190325

SCORE AT WEEK

W1

W2

W3

W4

W5

W6

(day 21)

Σ Random Forest - T21 - AUC 0.95

Follow-up queue

7

Search participants...

All

Needs attention

Check In soon

On track

P18

50%

Check in soon

Last active 20 days ago

P44

46%

Check in soon

Last active 19 days ago

P47

41%

Check in soon

Last active 19 days ago

P10

38%

Check in soon

Last active 5 days ago

P48

36%

Check in soon

Last active 18 days ago

P36

30%

Check in soon

Last active 12 days ago

P37

25%

Check in soon

Last active 16 days ago

P44

Cohort 1600 - 17500

Check in soon

<

>

Snooze 7d

Dismiss

46%

MEDIUM RISK

WHAT'S DRIVING P44'S RISK

Browsing programme content

Sparse facilitator presence in cohort

Few activities completed

↑ raises this person's risk · ↓ lowers it. Bar width = share of the explanation.

ACTIVATION LEVEL

Medium

Engagement is mixed — a focused nudge often helps.

RECOMMENDED APPROACH

→ Send a brief, personalised check-in message this week

→ Highlight an upcoming session relevant to their goals

WELLBEING SUPPORT CUE

Acknowledge contribution and invite a small next step without pressure.

ENGAGEMENT SIGNALS

COLLAPSE

LAST ACTIVE

19 days ago

DISCUSSION POSTS

0

vs prior 14d: n/a

FACILITATOR CONTACT

0

no comments yet

RECENT ACTIVITY

SHOW

Outreach · P44

P44hasn't posted yet

No activity in the current scoring window. For an at-risk participant, a warm first check-in is the highest-value early action — it's a dropout-lowering signal in the model.

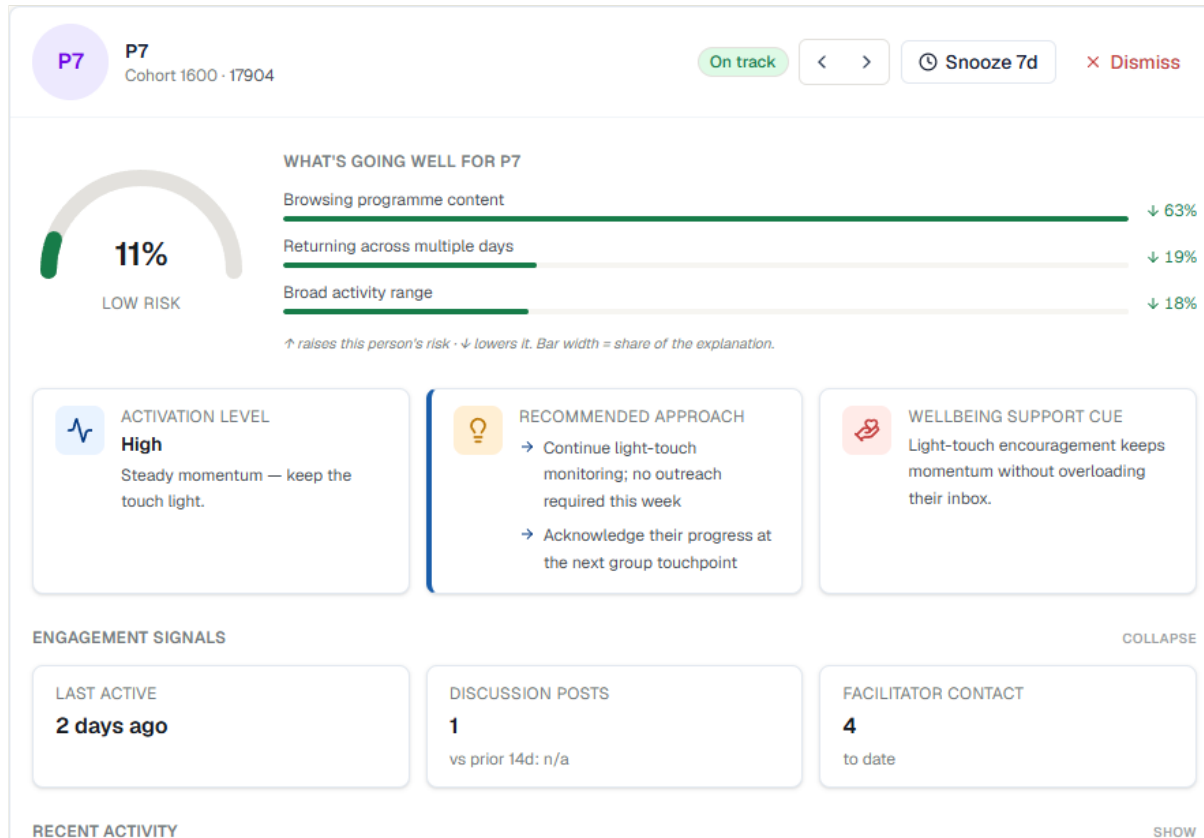
Write a first check-in

Generate drafts

Write my own

N

See who needs you and why



- Everyone **ranked each week**:
Needs attention · Check in soon · On track.
- A clear **risk score** with the **reasons** behind it.
- A **recommended approach** and wellbeing cue for outreach.

Risk comes from a **Random Forest** model, shown in plain language.

Reply with warmth, faster

Outreach · P7

Drafted by Qwen3 1.7B

PARTICIPANT POST

from 453d ago

GoalSetting

I am going to write a blog this week - topic TBC. on I will use my tablet . I will aim to do this Once let week

Regenerate

Write my own

SUGGESTED TONE

Warm check-in

Next-step nudge

Goal-focused

To: P7 · in-app message

Editable draft

P7, thank you so much for sharing your goal. We look forward to seeing all that you'll achieve in weeks ahead! 😊

113 chars

Helpful? 👍 🗨️ ⋮

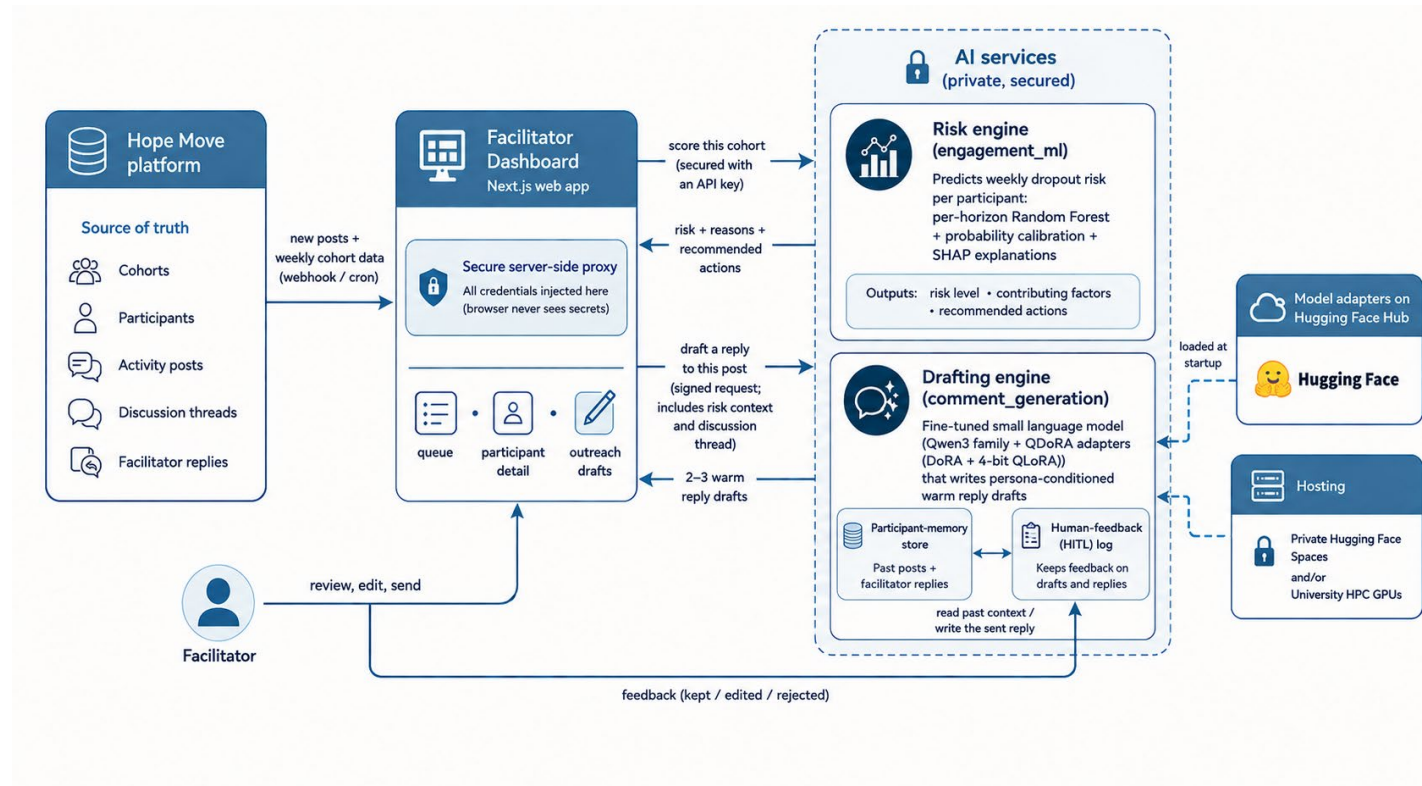
Send

What this draft is based on

- **Three tones**, drafted from the participant's own post.
- Written by a **fine-tuned small language model** (here, Qwen3 1.7b parameter model).
- **Edit, polish, send**: always your words, your call.
- Your **feedback** (kept / edited / rejected) trains better drafts.

The draft is a **starting point**, never sent on its own.

How it fits together



- **Risk engine:** predicts weekly dropout with a Random Forest, ~0.85–0.95 AUC.
- **Drafting engine:** a fine-tuned small language model trained on real Hope replies.

Both produce **suggestions only**; a **person always decides**.

THE RESEARCH BEHIND IT

Three research threads

1. Why retention happens · 2. Predicting dropout · 3. Generating replies

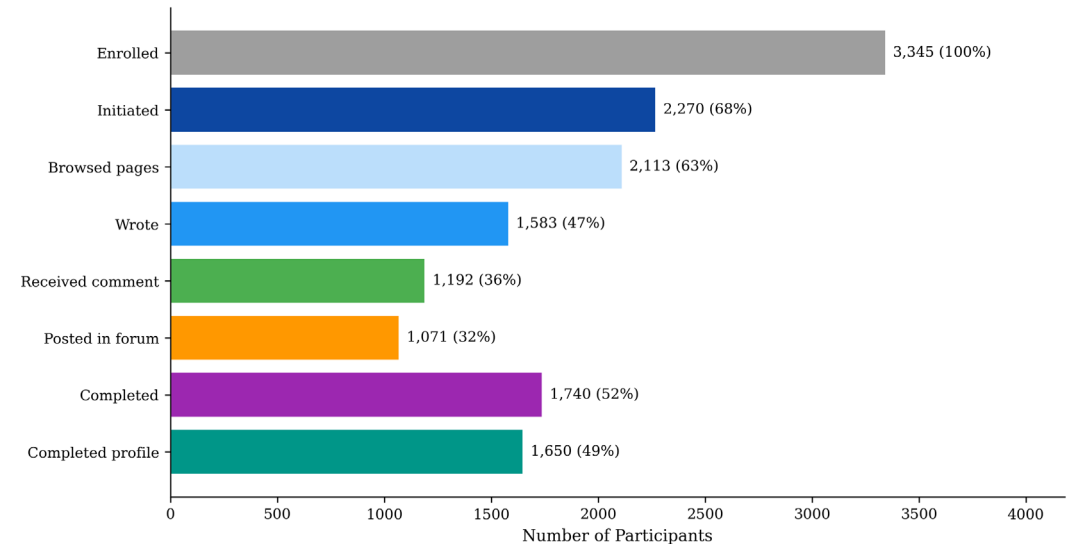
The data behind the research

Table 2: Study population by course. Enrolments includes all records; starters excludes non-initiators.

Course	Enrolments	Starters	Completion (%)	Writers (%)
NHS (2 modules)	1,300	911	79.5	72.8
Fire Fighters Charity	519	377	81.2	74.8
Macmillan Cancer Support	483	346	74.3	65.3
PCOS	328	229	72.9	72.5
Hope 4 The Community (MS)	290	201	69.2	60.7
IIH UK	216	100	74.0	58.0
MSK / HOPE MOVE	209	106	68.9	62.3
All	3,345	2,270	76.7	69.7

Study population by course: enrolments, starters, completion and writing

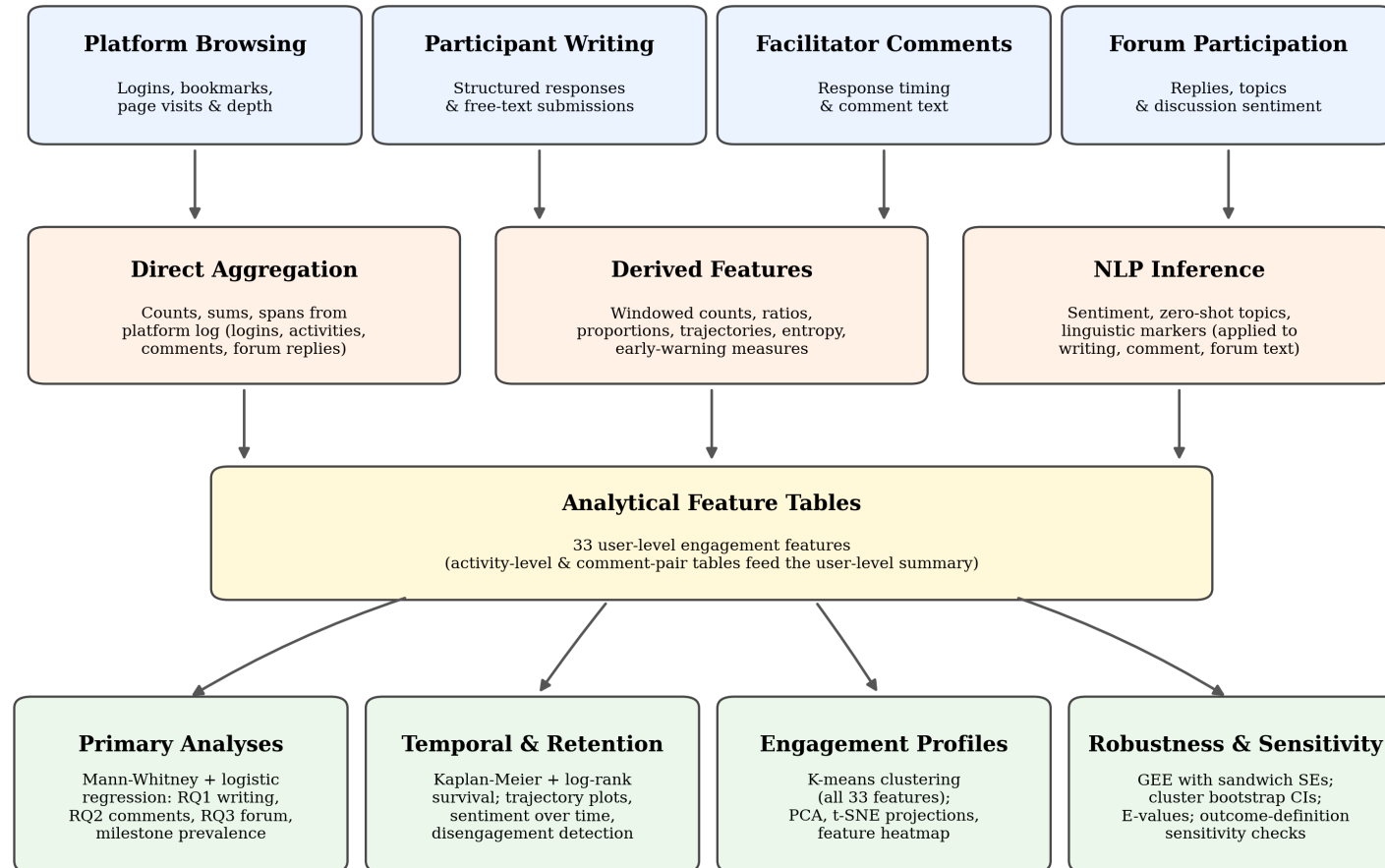
- **3,345 enrolments** across 7 courses (8 modules). Of these, **2,270 started**, and they form our analytic sample.
- **76.7% of starters completed**, and 69.7% wrote at least one activity.
- The biggest loss happens **before writing ever starts**.



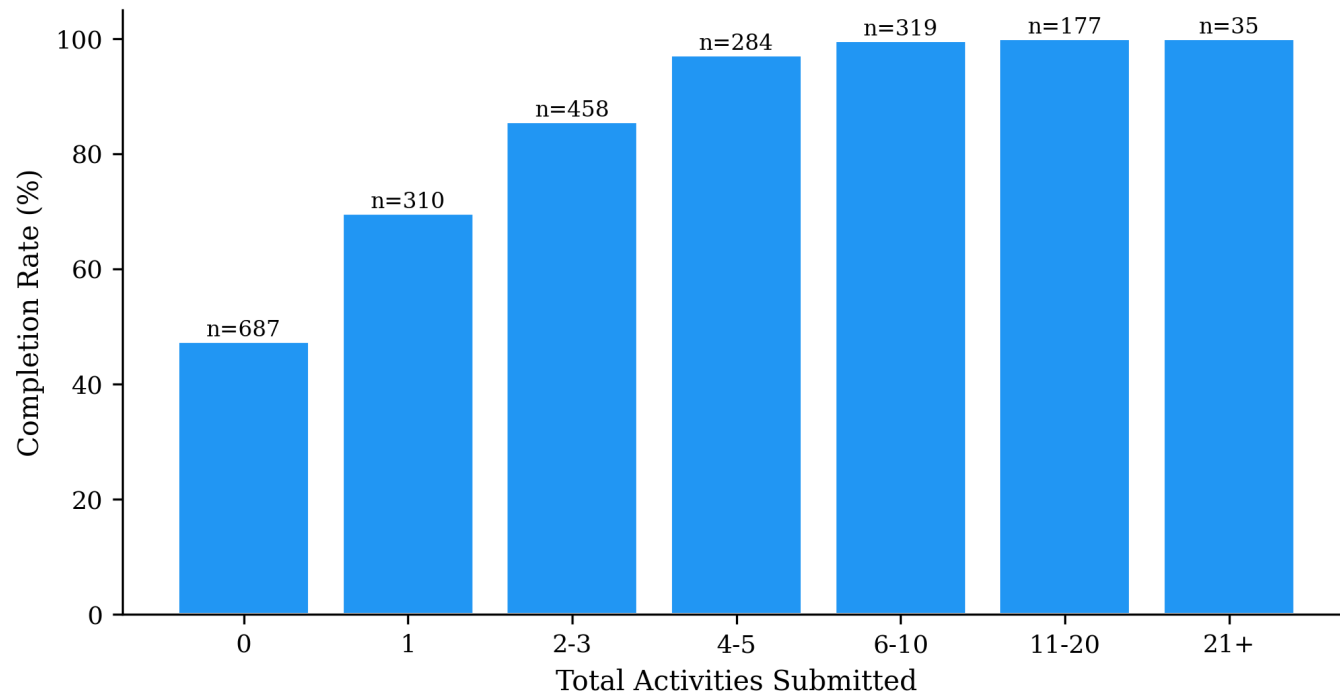
Engagement milestones across all enrolments (N = 3,345)

How we analysed engagement

- **Four data streams** from the platform: browsing, writing, facilitator comments, forum.
- **33 engagement features** per participant, including language signals from the text itself.
- **Four analysis streams:** group comparisons, retention curves, engagement profiles, robustness checks.

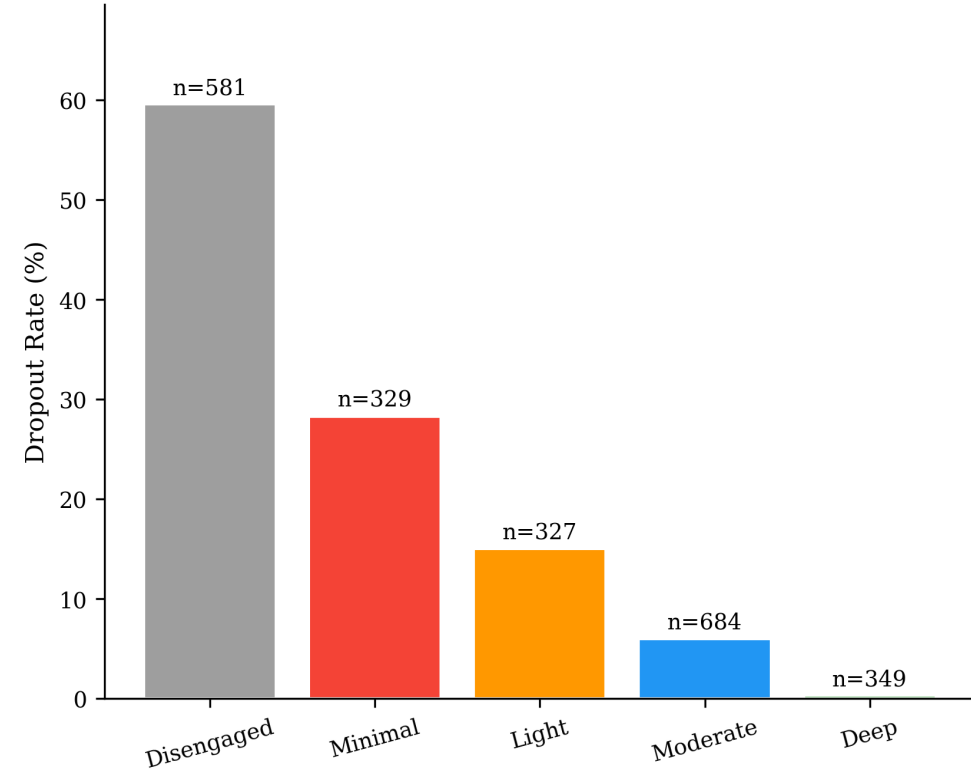
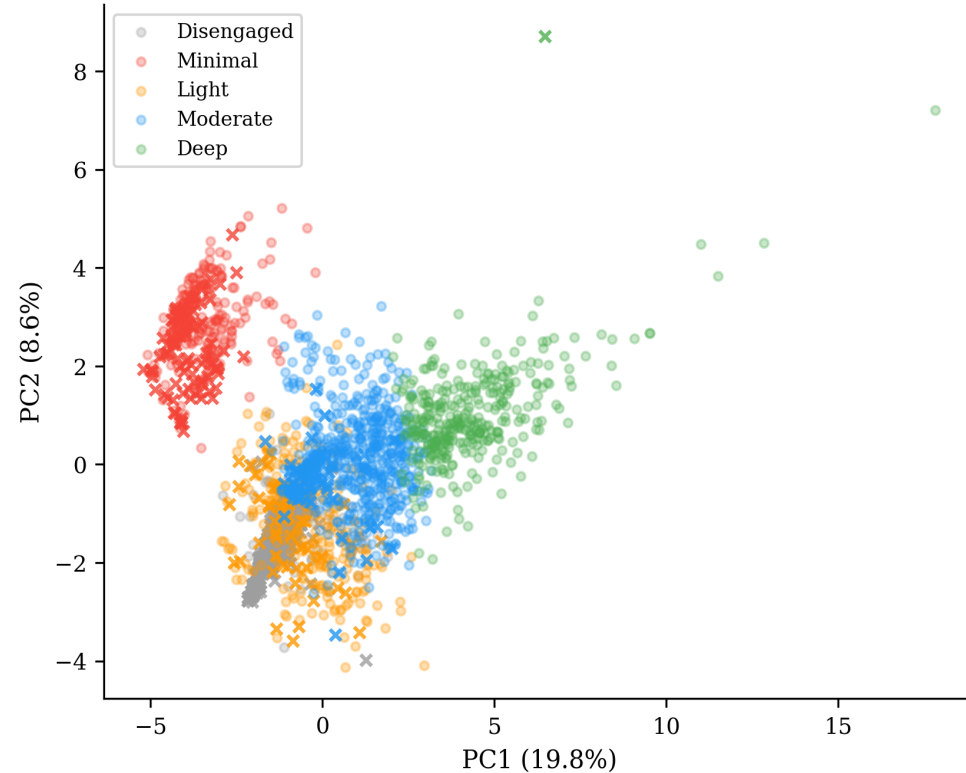


What keeps people engaged



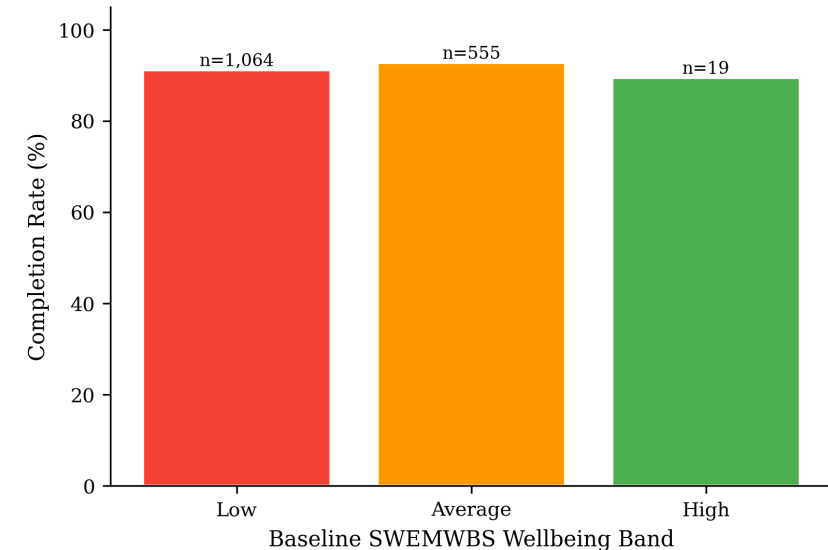
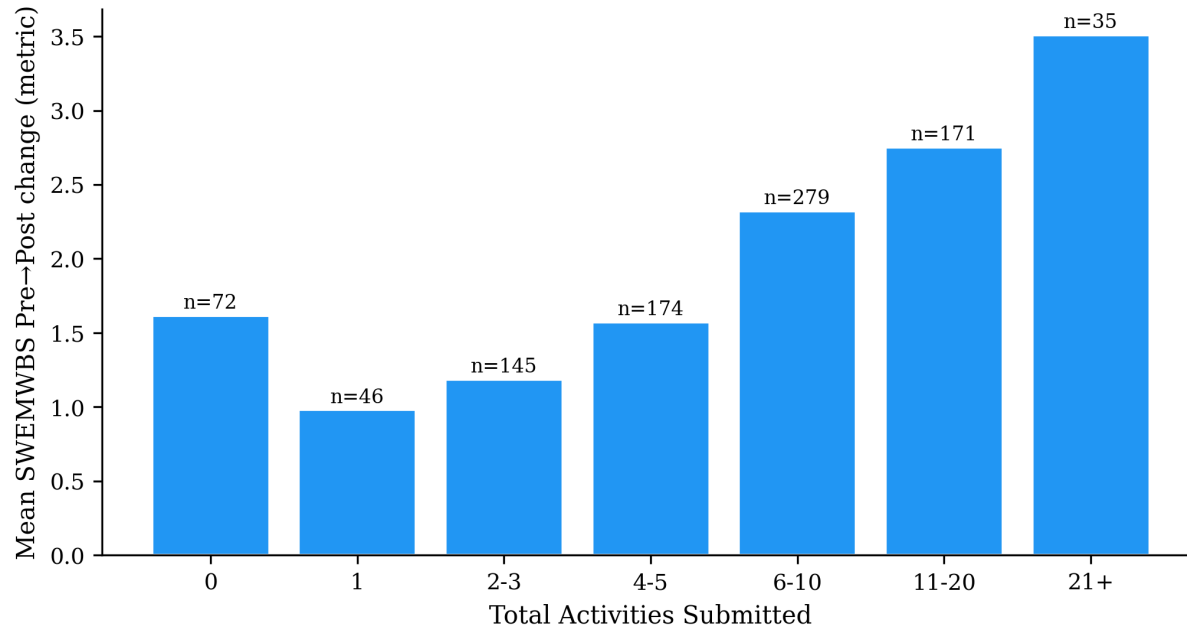
- **2,270 enrolments** across 7 courses (8 modules), delivered on the H4C digital platform.
- People who **write** complete at **89.3%**, against **47.5%** for those who don't. That is about **9 times the odds of completing** (OR 9.24, $p < .001$).
- **Dose-response**: completion climbs with every activity submitted, from 47.5% at zero to almost 100% from 11 activities.
- **646 enrolments** browsed the platform but never wrote. This is the group an early nudge could reach.

Five engagement profiles



- Drop-out ranges from **Disengaged (59.6%)** down to **Deep (0.3%)**.
- The '**Light**' group (n=327, 15.0% dropout) browses but barely writes. They are the **early-support target** the dashboard surfaces.
- Even a little writing matters. Dropout halves between Minimal (28.3%) and Light (15.0%), which makes the **first writing activity** the priority intervention.

Wellbeing rises with engagement

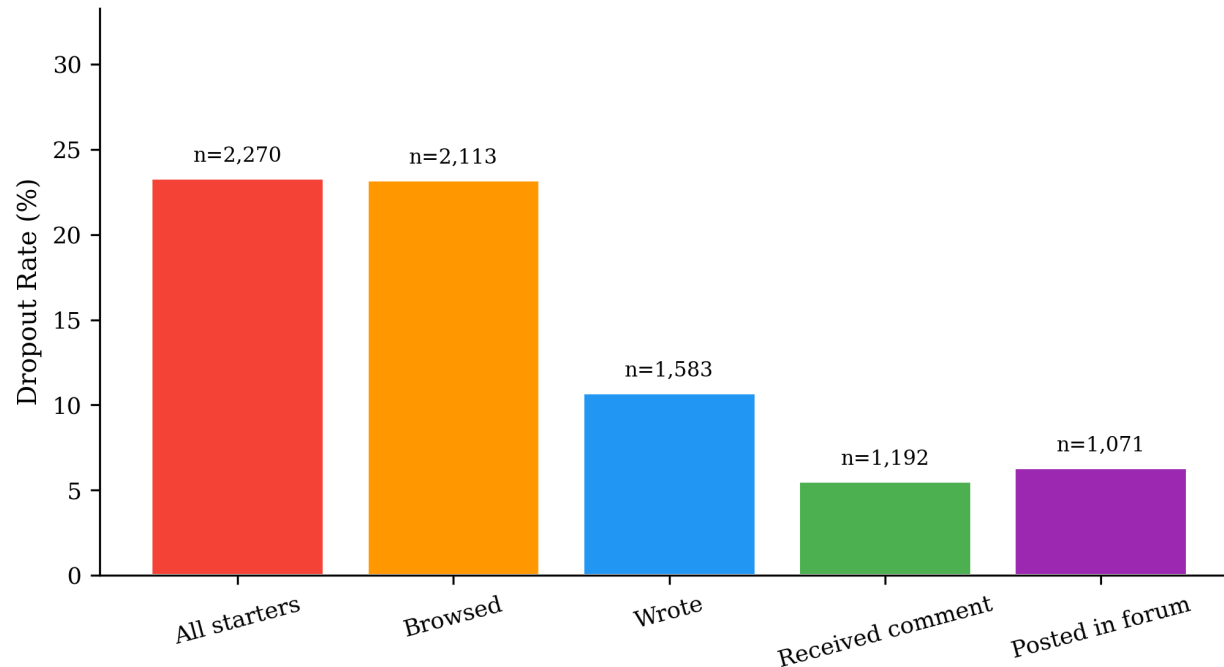


Completion is flat across baseline wellbeing bands

Mean wellbeing change from pre to post, by activities submitted

- **Median gain of +2 SWEMWBS points** from pre to post (922 paired responses, $p < .001$). More writing goes with larger gains, at **+0.10 points per activity** after adjustment.
- **Baseline wellbeing did not predict completion.** Early behaviour is the signal to watch, not how people felt at the start.

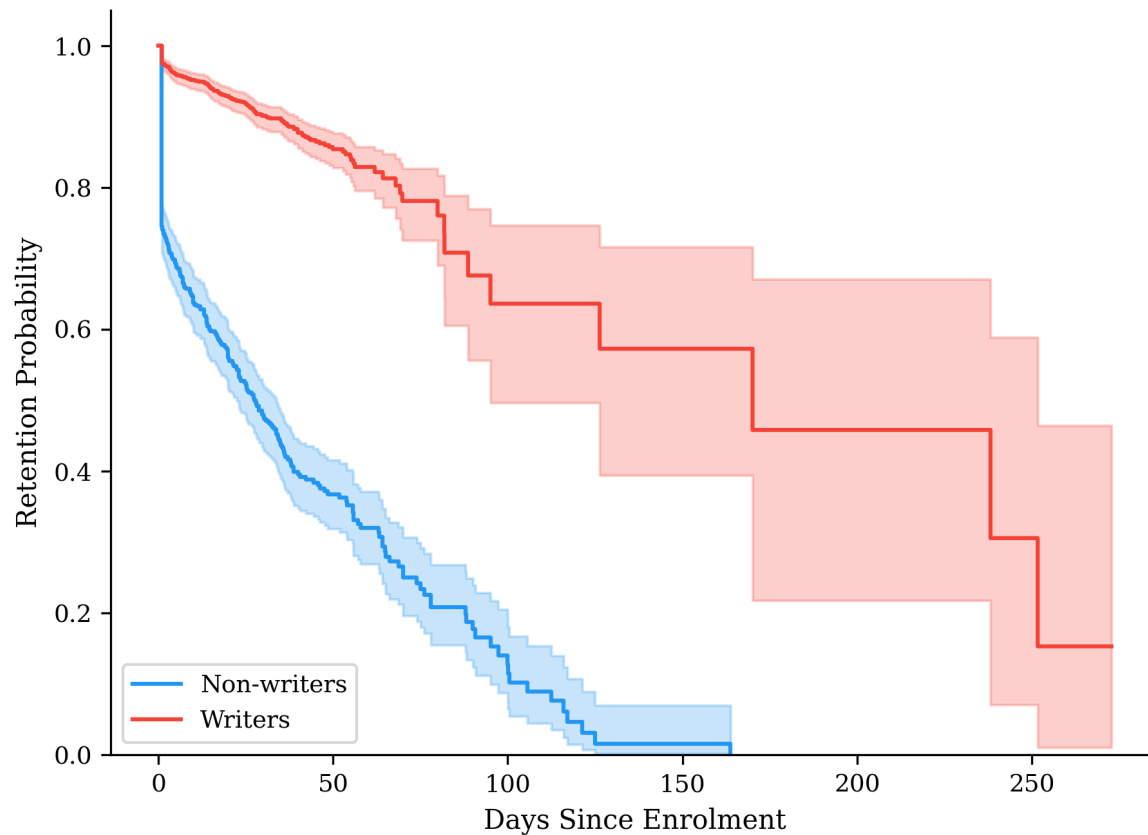
Connection lowers drop-out



Drop-out by the furthest milestone a person reached

- **Every milestone reached lowers drop-out:** 23.3% across all starters, **10.7%** once someone writes, and around 6% once they receive a comment or post in the forum.
- **57.6% of completers posted in the forum**, against 12.8% of those who dropped out.
- **Forum replies stay protective** after adjusting for writing, comments and logins (OR 0.78 per reply, 95% CI 0.68 to 0.89).
- **Posting early matters:** people who posted within 14 days completed at **93.0%**, against 61.5% for those who never posted.

The first week decides



Retention over time, writers (n = 1,583) against non-writers (n = 687)

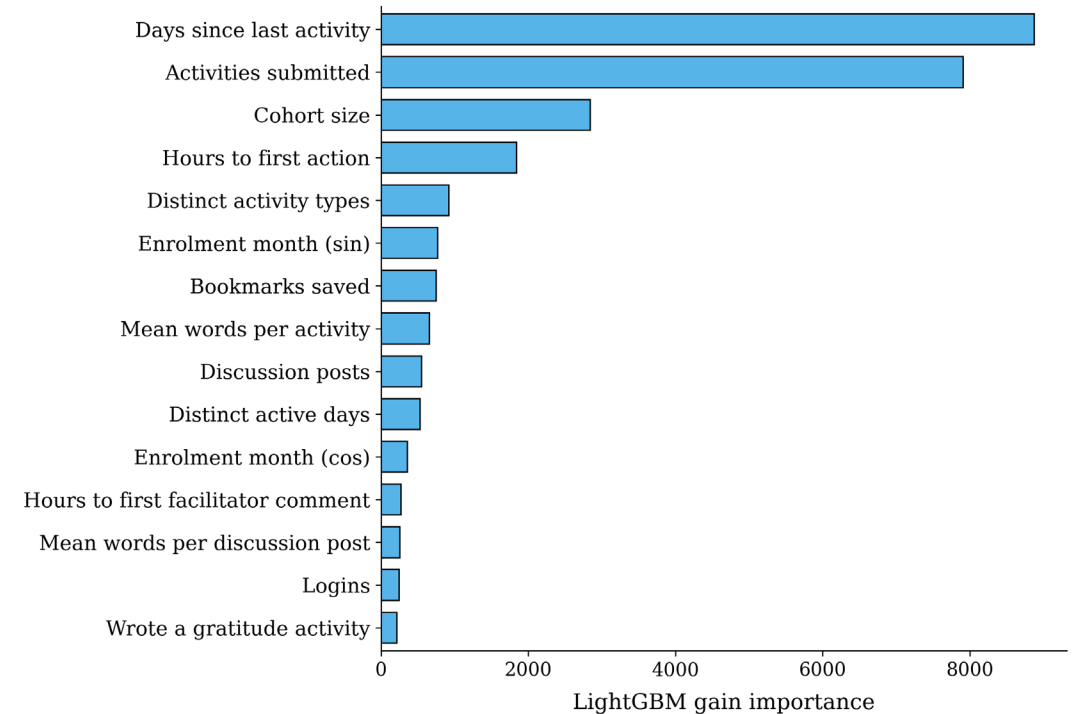
- **The two groups separate immediately** and never converge (log-rank $\chi^2 = 564$).
- Among writers who later dropped out, the **median last activity was day 0.7**. They arrived, wrote once, and stopped.
- A quarter of them left a **gap of more than 7 days** before disappearing, which is a visible warning.
- This is why the risk engine scores **every week, starting at day 7**.

Machine learning for the risk engine

Family	Models
Linear	Logistic regression
Tree (single)	Random forest
Tree (boosted)	LightGBM, XGBoost, CatBoost
Shallow NN	MLP
Recurrent	GRU

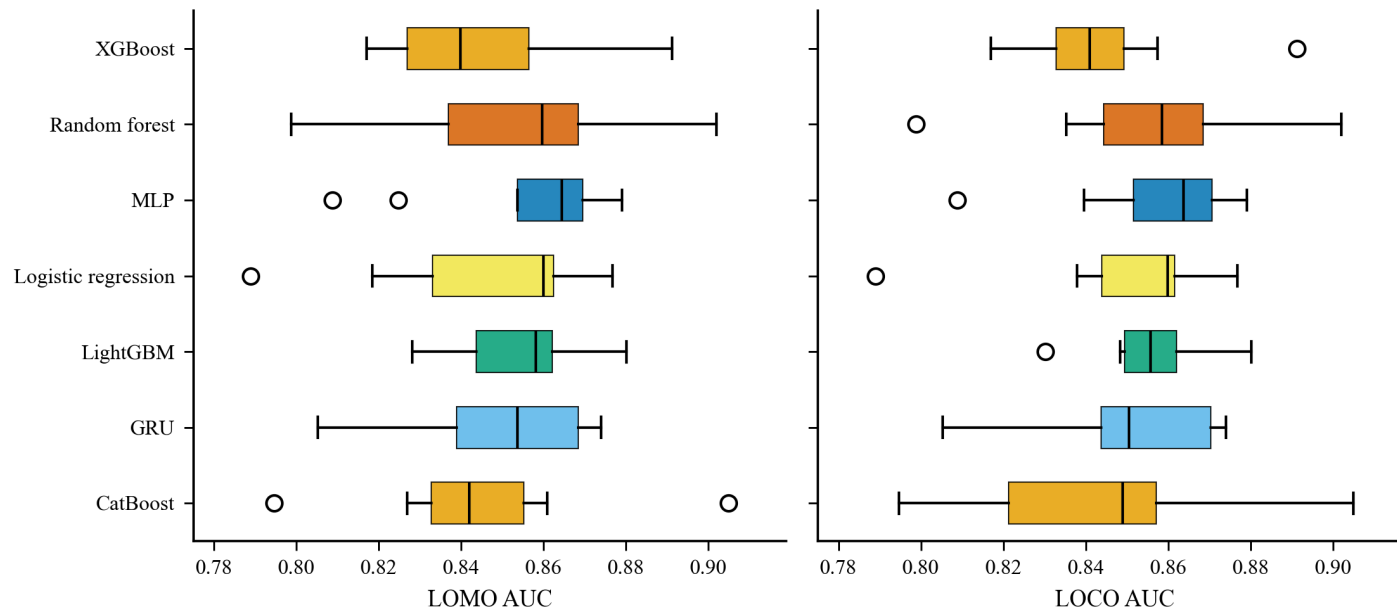
Weekly checkpoints at $T = 7$ to 56 days. Each model sees only the first T days of behaviour.

Top predictors: days since last activity, activities submitted, hours to first action.



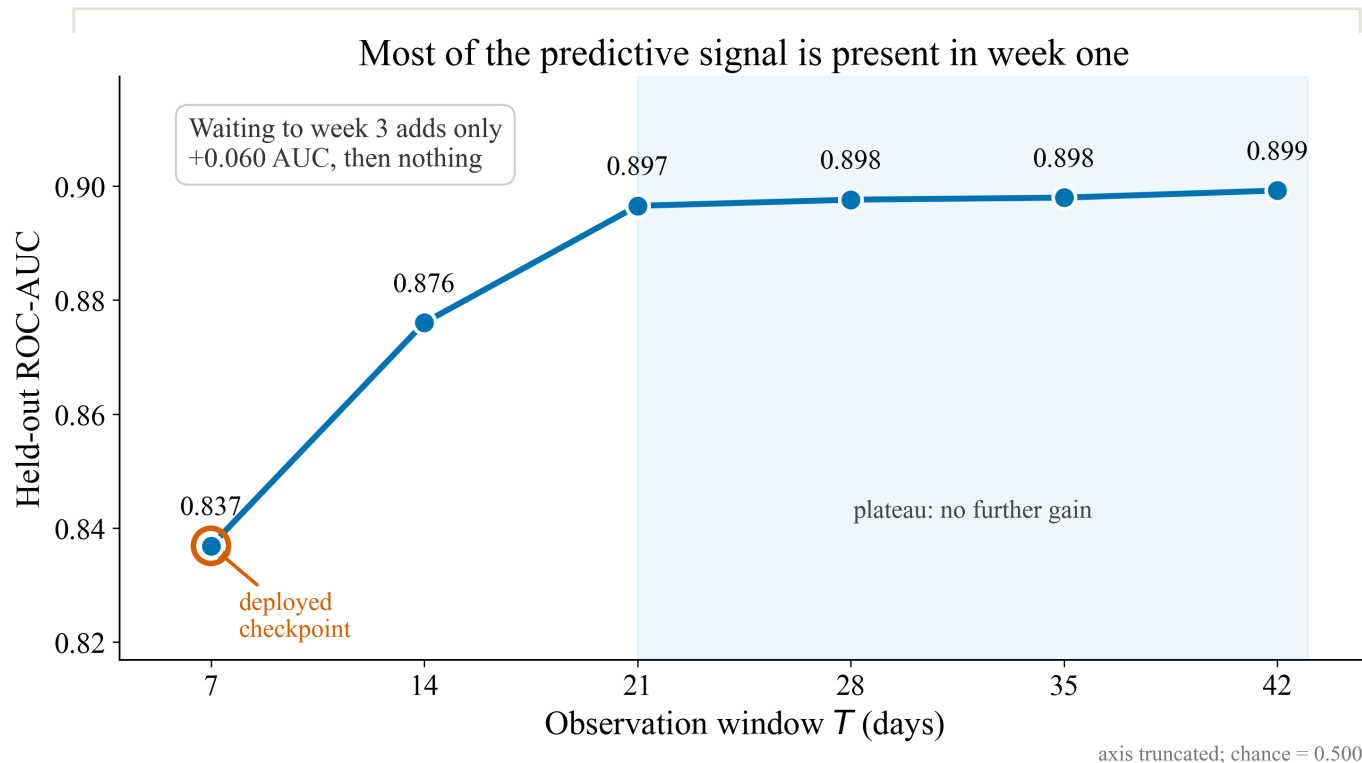
What the deployed model relies on (gain importance)

Choosing the model



- **Seven model families** compared (LR, RF, XGBoost, LightGBM, CatBoost, GRU, MLP).
- Tested on **held-out whole courses (LOCO) and modules (LOMO)** the model never saw.
- The top models are **statistically tied** at about 0.85 AUC, with no family separating from the pack.
- **LightGBM is deployed** for engineering reasons, namely calibrated scores and built-in explanations, rather than because it won.

Predicting dropout early



- **AUC 0.84 at week 1**, rising to **0.90** by course end, so dropout is visible early.
- One model per weekly checkpoint, and **most of the signal is already there in week one**.
- Transfers to **courses and modules the model has never seen** (about 0.85 AUC held out).
- Tuned for **recall of 90%**: at week 1 it catches **601 of 667** eventual dropouts, at a precision of 0.72.

Breadth beats depth

94.5%

writers who received a facilitator comment
completed

73.4%

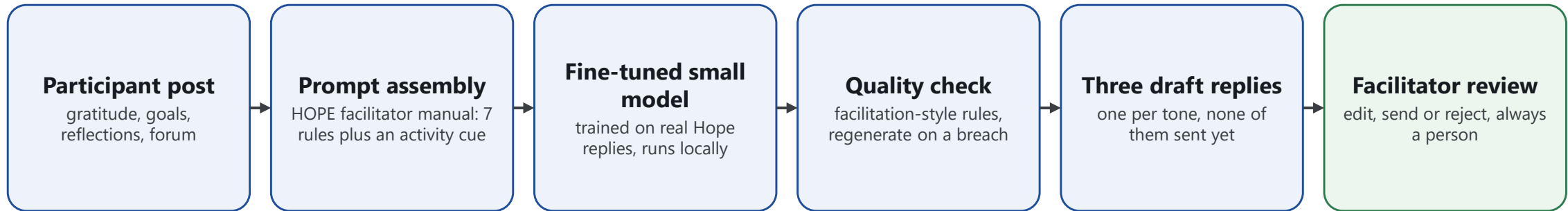
writers who received none completed

no effect

of how fast, or how long, the reply was

- **Being contacted at all is what tracks with completion.** Response speed ($p = 0.56$) and comment length ($p = 0.83$) made no measurable difference.
- Once early writing and logins are accounted for, the effect is **no longer distinguishable** (adjusted OR 0.79, $p = 0.32$), so we read this as a marker rather than proof of causation.
- The practical reading: **reach more people with a short warm reply** rather than perfecting a few. Facilitator time is the scarce resource, and that is exactly what the drafting engine is for.

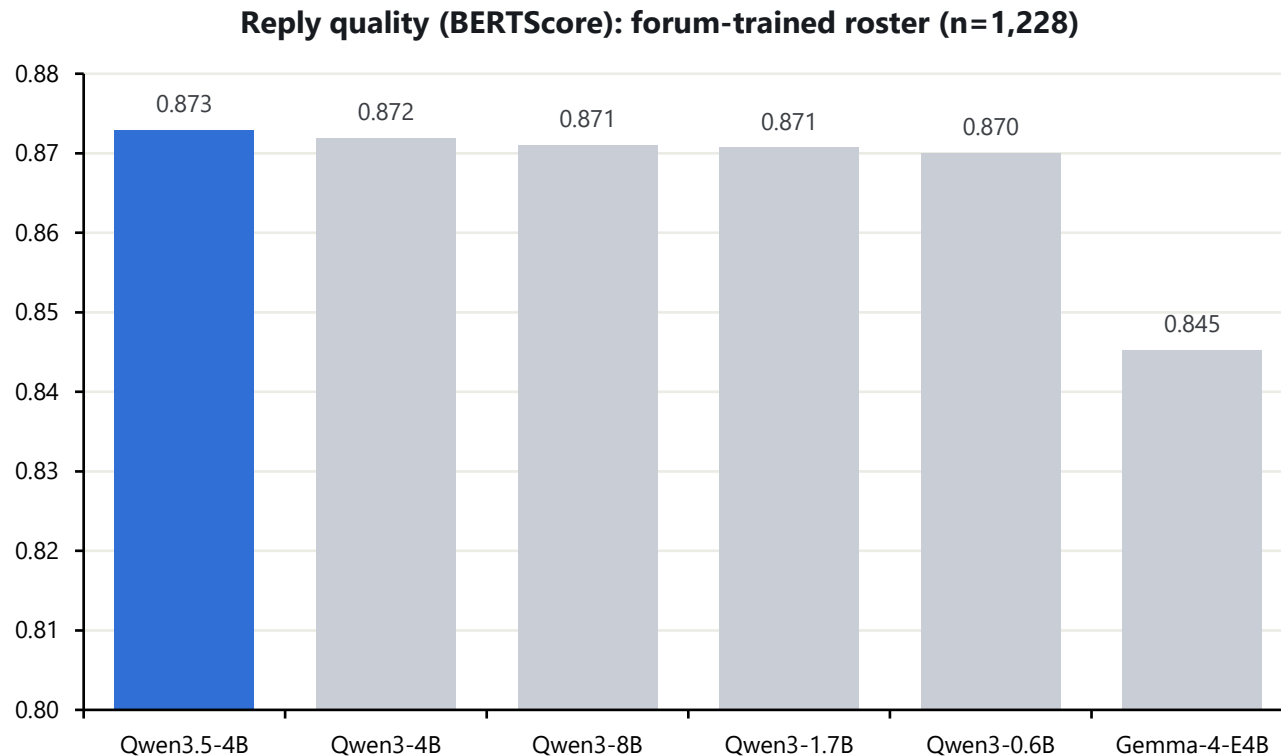
How a draft is made, safely



kept, edited or rejected: that feedback improves future drafts

- Trained on **3,900 cleaned pairs of real posts and facilitator replies**, plus **2,313 forum threads**, written by 46 Hope facilitators.
- **The manual is the rulebook**: reflect back, ask one open question, never advise or diagnose.
- Everything runs on our own hardware, so **no participant text leaves the machine**.

Drafting replies that sound like Hope



- **8 open models** fine-tuned with **QDoRA** on more than **3,900 real pairs of Hope posts and replies**, using programme data only with no synthetic augmentation.
- **Qwen3.5-4B** is the deployed default, with the top BERTScore of 0.873.
- **Qwen3-4B** ties it within noise, 0.001 apart, and has the **best goal-setting fidelity**.
- **Safety held across every model**, with no medical advice, no diagnosis and no links. A person always sends.

Cleaning mattered. An audit caught drafts congratulating goals that had not been attempted yet, and retraining on cleaned pairs cut this from about 30% to 1% of drafts, re-ordering the model ranking.

Three things a programme can act on

1

Get the first activity written

646 people browsed but never wrote. That first piece of writing is the single biggest lever we found, and it is the moment to spend effort on.

2

Get people into a conversation early

Posting in the forum within 14 days tracks with 93.0% completion. Prompting an early reply is cheap and it is protective.

3

Spread facilitator time wider

Presence matters more than polish. Reaching everyone briefly beats writing a perfect reply to a few.

The AI exists to make those three things possible at cohort scale, **with a facilitator deciding every time.**

Honest science, careful rollout

What we're careful about

- **Observational data**, so writing may signal commitment rather than cause completion.
- **One UK platform and provider**, so findings need revalidation elsewhere.
- **Week-1 risk scores rank well** but are not yet calibrated probabilities.
- **No demographics in the data**, so a fairness audit is a precondition for clinical use.
- **The dashboard holds no models and no participant database**, and credentials stay server-side.

Where this goes next

- **Learn from facilitator edits**, so that kept, edited and rejected drafts become training signal.
- **Validate on other programmes** and new cohorts, which is the claim most in need of replication.
- **Agree alert thresholds with each service**, since the balance between catching people and false alarms is theirs to set.

Nothing reaches a participant without a facilitator deciding to send it.

Thank you

HOPE-MOVE: AI to help facilitators keep people engaged and supported.

Matthew England: Matthew.England@coventry.ac.uk

Michael Ajao-Olarinoye: ae6780@coventry.ac.uk

