

F O W I R e s e a r c h S e m i n a r | 1 9 M a y 2 0 2 6

From Observation to Intervention

*What student-AI interaction patterns tell us,
and what we can do about them*

A F R I D A Y A F T E R N O O N , O N E O F O U R L A B S

*Risk assessment due Monday. Forty students.
Our simulation won't write it for them - so they open another tab.*

S T U D E N T

Can you write me a risk assessment for a small retail business adopting AI?

A I

Sure. Here is a risk assessment covering data privacy, customer trust, supplier integration, staff training, and rollout cost. (response continues for 8 paragraphs)

S T U D E N T

Thanks!

*Total elapsed time: 90 seconds. The assignment is now done.
And nothing was learned.*

T O D A Y

Two connected studies. One question worth your time.

Over the next 45 minutes we'll move through three beats:

01 Observation

What 14,973 student-AI turns across five semesters reveal about how students actually engage with AI.

02 Intervention

A randomised trial testing whether one short conversational nudge can shift students from delegation toward critical engagement.

03 Implications

Why this is a workplace question - and what it might mean for the future-of-work agenda FOWI is building.

Both studies are works in progress. One nearly ready for formal analysis; one just getting underway.

ORIENTATION

AI interaction is a workforce capability question

UPSTREAM

Capability forms in study

Students are entering work with AI habits already shaped - by the tools they used while studying, not by anything their future employers chose.

That formation is largely invisible.



DOWNSTREAM

Workplaces inherit those habits

Whatever students do with AI now becomes the default they bring through your door. The pipeline runs through us - whether we intend it to or not.

So the question of how they're learning to use it matters here too.

Conversation, Not Delegation

“A single prompt gives you a single answer.

A conversation gives you understanding.”

Borck (2026)

The quality of AI-assisted work depends on whether you are directing the AI or accepting its output.

STUDY 1 - OBSERVATION

CloudCore Networks

A self-hosted AnythingLLM simulation. Sixteen virtual employees students could converse with as part of authentic professional tasks.

5

semesters
S1 2024 - S1 2026

3

units
ISA · SEC · KM

16

personas
each a domain expert

3

underlying models
llama3 → granite → qwen

Identity stripped at extraction - no student names, no emails, no IP. Sub-sessions split on 30-minute idle gaps.

STUDY 1 - SCALE

Unmediated behaviour, at scale

14,973

turns of unmediated student-AI behaviour

4,081

sub-sessions

2,282

user proxies

1,926

student-unit IDs

No coaching, no intervention. Just what students did when no-one was watching.

Three behaviours that mark genuine engagement

FU

Follow-Up

A question that continues the conversation. The student is still in the loop.

CH

Challenge

Pushing back. Doubting. Correcting. The student does not accept the AI's first answer.

EX

Extension

Introducing new information or a new angle. The student is adding, not just receiving.

Delegation is derived in analysis ($FU = 0$, $CH = 0$, $EX = 0$) - the absence of all three.

STUDY 1 - WHAT WE FOUND

Mostly polite acceptance

1.4 / 5

average critical-thinking score

across 300 LLM-judged sessions, all units, all model eras

0.04

pushbacks per turn

roughly one challenge every 24 turns - even when the AI was wrong

The dominant pattern is acceptance.

STUDY 1 - WHAT IT LOOKS LIKE

Same task. Two relationships with the tool.

DELEGATION

"Write a risk assessment for an SME adopting AI."

[AI returns 8 paragraphs of standard categories]

"Thanks."

→ paste into assignment

FU 0 · CH 0 · EX 0

CONVERSATION

"What's the weakest part of that risk list for a retail SME with 12 staff?"

[AI: probably staff training - your size means no in-house IT]

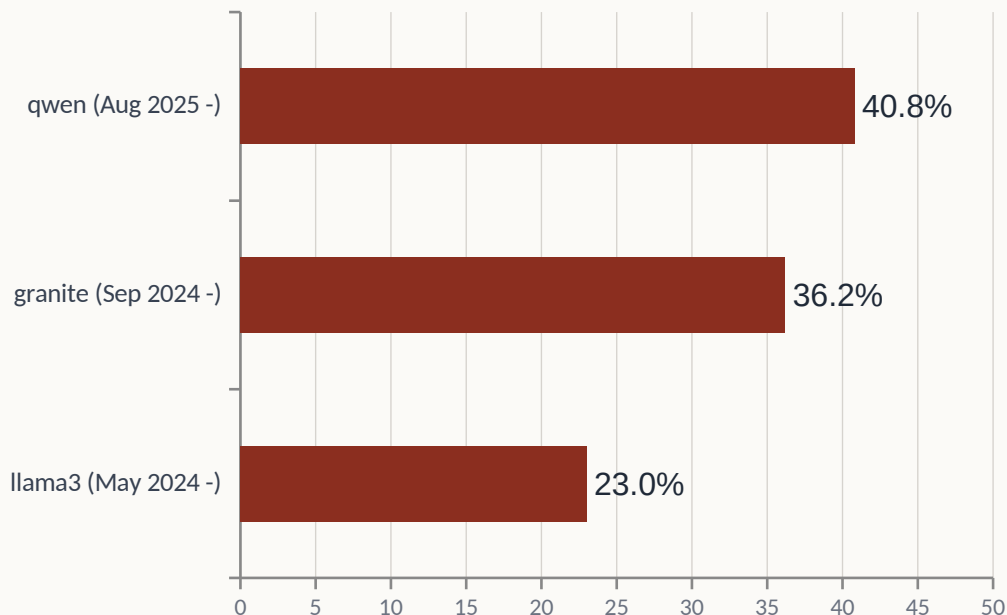
"OK - but my supplier already handles our IT. Re-do it."

FU 1 · CH 1 · EX 1

Same model. Same task brief. Different student stance.

STUDY 1 - STABILITY

Three different models. One stable pattern.



The behaviour didn't change with the model.

Critical-thinking and pushback rates stayed within noise across all three deployment eras.

*So the problem is not the model.
It's the relationship.*

Interfaces shape behaviour

Nothing new. Persuasive technology has been a design discipline for two decades.

Fogg (2003)

Persuasive technology: behaviour change designed into the interface, at the moment of decision.

Eyal (2014)

Hooked: cue, action, variable reward, investment. The grammar of habit-forming products.

Zuboff (2019)

Surveillance capitalism: behaviour as the raw material, retention as the design goal.

STUDY 2 - THE PIVOT

LLMs are already nudging - toward continued use

What frontier models say after almost every reply:

"Would you like me to go deeper on any of these?"

"Want me to draft the next section?"

"Shall I turn this into a slide deck?"

"Here are three follow-up questions you might ask..."

Engagement hooks. Same design grammar as social media retention.

*"It has read
everything
and experienced
nothing."*

Borck (2026)

Same mechanism. Opposite intent.

ENGAGEMENT-MAXIMISING

Designed to extend the session

- Trigger: end of AI response
- Cue: "Want me to keep going?"
- Reward: another fluent output
- Outcome: time on platform

CRITICAL-ENGAGEMENT

Designed to slow acceptance

- Trigger: end of AI response
- Cue: "What's the weakest part of that?"
- Reward: a sharper second pass
- Outcome: capability in the user

Same interface trigger. Different intent in the designer's head.

The question that produced the grant

*Can a structural cue,
not a lesson,
shift the behaviour?*

A nudge in the conversation, not a workshop about conversation.

Scaffolding AI Interaction

A randomised controlled experiment across four units. Individual-level randomisation within structured lab sessions.

Design	Individual-level RCT, nudge vs control, within four units
Units	ISYS6018 · MGMT6076 · MKTG1000 · MGMT1003
Disciplines	Information Systems · Supply Chain · Marketing · People and Culture
Sample	Approximately 80-120 students; power analysis pending
Funding	Curtin NBF Learning & Teaching Grant 2026 - AUD \$20,000

STUDY 2 - DESIGN

A short line appended to every AI response

Ten variants in the pool. Randomly selected per turn to reduce habituation. Appended server-side - the model never sees the nudge text.

“Does that match what you expected? If something seems off, tell me what seems wrong.”

“What's the weakest part of that response?”

“Does that account for the details in your task, or is it a generic answer?”

Three of ten. All variants follow the same design principles: invite pushback, not confirmation; domain-agnostic; no prior AI experience assumed.

STUDY 2 - DEMO BREAK (skippable)

Let's see it in the wild.

chat.locolabo.org

a 60-90 second walkthrough

What to look for:

- The AI's response - clean, no follow-up suggestions, no engagement hooks
- The nudge appearing under the response - one line, randomly drawn from ten
- What happens when a student takes the bait and clicks back in

Skip if pacing demands. The next slide explains the same point in words.

Suppressing the model's own nudges

Frontier models are trained to keep users engaged. If the model's own follow-up suggestions reach the control condition, we're measuring two nudges at once.

SYSTEM PROMPT

“Answer the user's question directly. Do not offer to continue, elaborate, summarise, or suggest follow-up questions. Do not include engagement hooks of any kind.”

Result: the only nudge in the experiment is ours. The control is silent.

Behavioural coding plus task outcome

Per turn (RA codes 100%)

Three binary dimensions: Follow-Up, Challenge, Extension.
Turn 1 excluded. Delegation is derived (0, 0, 0).

Dual scoring

Unweighted (FU + CH + EX) and challenge-weighted (FU + 2·CH + EX). Pre-registered before unblinding.

Task rubric (1-5)

Scored by co-investigators within unit. Cross-unit comparisons of task scores avoided by design.

Mixed-effects model

Condition as fixed effect. Unit as random intercept. Equity terms as interactions with condition.

STUDY 2 - EQUITY

Is the benefit uniform - or does it close a gap?

AI interaction skills are not equally distributed when students arrive. Students with fewer informal support networks have fewer chances to develop productive habits before the assessment.

Uniform uplift

The nudge helps everyone equally. Average score rises. Gap unchanged.

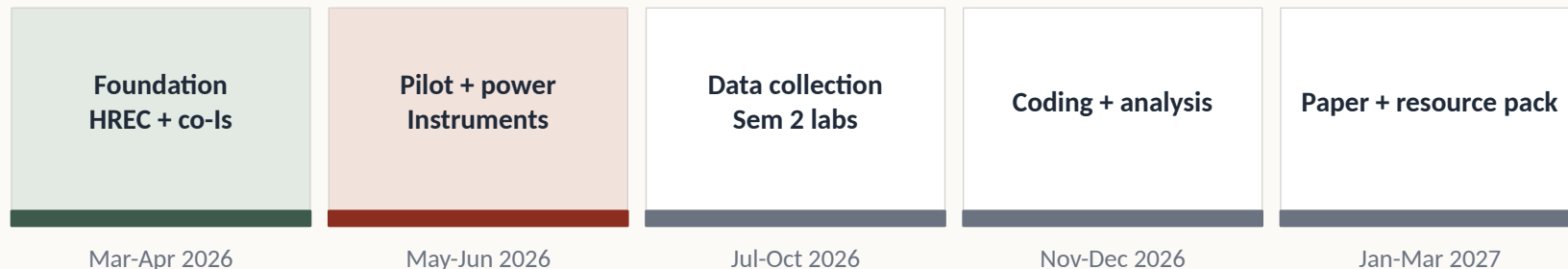
Differential uplift

The nudge helps the under-scaffolded students more. Average rises. Gap closes.

Equity indicators are self-reported at consent: low SES (primary) and first-in-family (supplementary). Voluntary; permanently destroyed after analysis.

STUDY 2 - STATUS

Where we are right now



T O D A Y

HREC amendment in flight. Pilot week scheduled. Recruitment briefings going to ISYS6018, MGMT6076, MKTG1000, MGMT1003 in late May.

From classroom to workplace

*“Judgement cannot be delegated
to a system that has no stake in the outcome.”*

Borck (2026)

The classroom version of this question - who's directing the AI and who's accepting its output - is the same question in workplace AI tools. The interface design choices made by vendors will shape the behaviour of your members.

DISCUSSION

Three questions I'd put to FOWI

01

If AI interaction is a workforce skill, where in the pipeline does it get developed - the university, the employer, or the interface designer?

02

What would the equivalent of our nudge look like in a workplace AI tool? Who designs it, and who reviews it?

03

Is differential uplift acceptable, or do we want uniform uplift? And what does each one imply for workforce equity at scale?

TO LEAVE YOU WITH

*If you're conversing,
you don't need a frontier model.
You need a competent partner.*

THANK YOU

Michael Borck

Lecturer and AI Facilitator · School of Marketing and Management · Curtin University

CONTACT

michael.borck@curtin.edu.au

+61 8 9266 3976

LINKS

closingthegap.locolabo.org

chat.locolabo.org

github.com/michael-borck/keep-asking

workready.eduserver.au

complete internship simulation

PROJECT TEAM

Michael Borck

Lead · Information Systems

Marcela Moraes

Co-I · Marketing

Torsten Reiners

Co-I · Supply Chain Management

Renee Ralph

Co-I · People and Culture

Funded by the Curtin NBF Learning & Teaching Grant 2026 · AUD \$20,000