



Business Management Institute

Would You Let It Sign Off?

AI judgment and the decision that stays yours

Stefani Langehennig, PhD
Daniels College of Business
University of Denver
August 2026

CHECK-IN POINT

- Point your phone's camera at the QR code
- or*
- Navigate to <https://wacubo.cnf.io>
- Tap session matching this room
- Keep browser window open to check out at session's end

WACUBO 2025

Business Management Institute





Agenda

A practical framework for where AI belongs and where it does not.

- 01 The judgment gap** Why use and delegation are different
- 02 What AI is** A plain-language mental model
- 03 Where AI helps** Augmentation, not automation
- 04 Where risk enters** Errors, context, bias, accountability
- 05 A leadership framework** Low consequence vs. high consequence



The Judgment Gap

The contrast between use and delegation





Two questions for the room

?

How many of you have used
ChatGPT – or any AI tool – in the last
month?

?

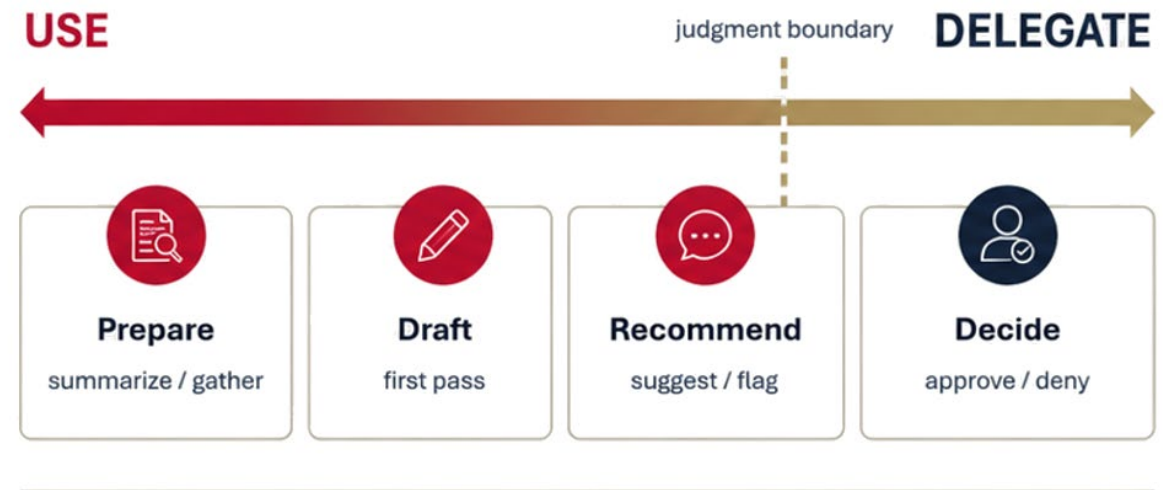
How many of you would let it make a
budget decision?

That gap between the two answers is what this session is about.

The question isn't whether AI shows up

It already has! The leadership task is boundary-setting.

- Staff are experimenting with AI whether adoption is official.
- Using AI to prepare work is different from letting AI make consequential decisions.
- Leaders need explicit expectations: where AI is acceptable, where review is required, and who is accountable.



Human review required before AI crosses into consequential decisions.



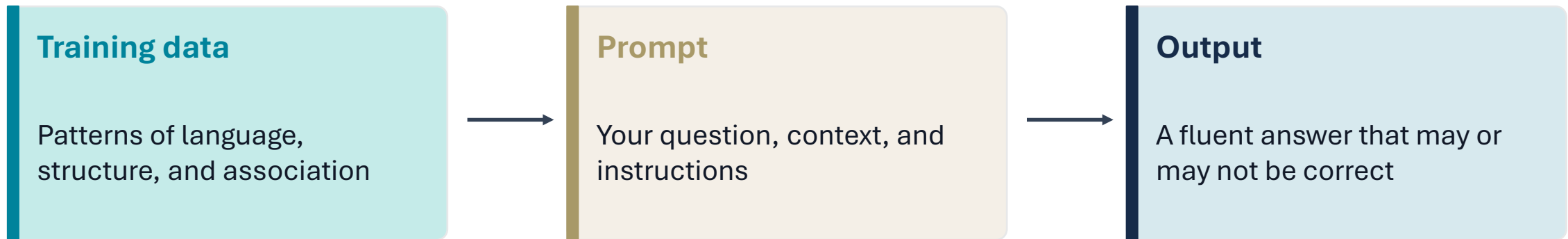
What AI Actually Is

A non-technical mental model



AI is a prediction machine

It generates plausible next outputs based on patterns.



The model does not “know” whether the output is true. It knows what answer looks plausible.



Fluency isn't the same as reliability

This is the trap for experienced leaders.

Human advisor

Hedges, uncertainty, and caveats may carry real information about what the person knows.

AI output

Confident tone is a language pattern. It is not proof of understanding, confidence, or correctness.

A polished answer can still be a wrong answer!

The plausible answer problem

Plausible is often useful, but plausible without verification is dangerous.

Low stakes

Draft a meeting agenda

Error is easy to catch and cheap to correct

High stakes

Summarize a travel policy for enforcement

Error can create compliance, budget, or trust problems

Leadership rule: plausible must become verified before it becomes institutional.



60 seconds: spot the problem

Compare the AI-drafted summary to the actual policy. Raise your hand when you find the error.

YOUR ACTUAL POLICY (EXCERPT)

Meal per diem is \$75 per day for domestic travel.

Receipts are required for all expenses over \$25.

Reimbursement requests are due within 60 days of return.

THE AI-DRAFTED SUMMARY

“Travelers may claim a \$75 daily meal per diem for domestic trips. Receipts are required for all expenses over \$25. Reimbursement requests must be submitted within 90 days of return.”

Two of three details are right. This is what validation actually looks like. Now, imagine 40 pages of this!



Where AI Helps Leaders

Augmentation reduces cognitive prep work so humans can focus on judgment





Think augmentation, not automation

Ai is strongest when it helps prepare work for human judgment.

Summarize

First-pass distillation of long documents

Draft

Starting points for communications

Explore

Options, scenarios, alternatives

Translate

Finance language into campus language

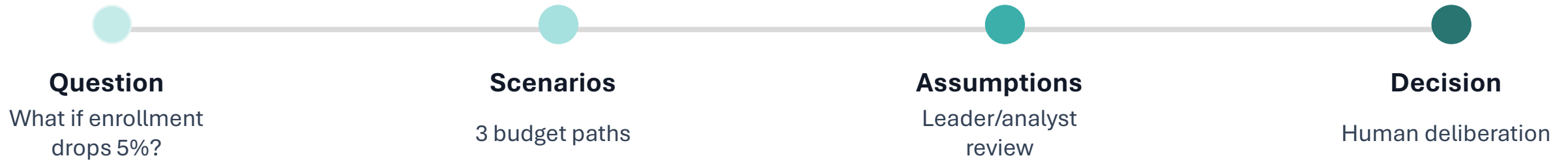
Organize

Turn messy inputs into structured outlines

AI does prep work. Leaders still evaluate, decide, and take responsibility.

Example: budget scenario exploration

AI can structure the scenario, but it can't validate the assumptions.



Helpful prompt approach: “Generate options I can evaluate,” not “Tell me what to do.”



Where AI Creates Risk

Leaders must employ disciplined skepticism



Risk map for administrative leaders

The risks are both technical and organizational.

Hallucinations

Fabricated citations, statistics, policy provisions, or facts

Missing context

Board politics, labor agreements, state formulas, local history

Bias

Existing inequities hidden inside language or modeling assumptions

Accountability gaps

Unclear responsibility for AI-assisted decisions

Vendor claims

Overstated automation and unclear data practices



Hallucinations are confidence without reality

AI can fabricate information in a voice that sounds institutionally fluent!

A fabricated accreditation standard or federal regulation can look entirely legitimate to someone who does not already know the correct answer.

Higher ed version: the fake answer sounds like something your institution would write.

Missing context: AI doesn't know your institution

Leaders make decisions inside local history, constraints, and relationships.

board politics

Enrollment
trajectory

Labor
agreements

Org-chart
history

Local exceptions

Campus climate

Donor
commitments

State funding

Every AI output is context-free unless context is provided, and even then it may not use it correctly.

Bias can hide behind objectivity

AI-assisted analysis can reproduce historical patterns under a neutral-looking guise.



- Hiring language and candidate screening can import bias.
- Student outcome models can reflect past inequities.
- Budget modeling may perpetuate historical funding patterns.

Accountability gets muddy fast

Who is responsible when an AI-assisted recommendation gets it wrong?

Staff member

Prompted or used the tool

Leader

Approved the work

IT/CIO

Approved the system

Vendor

Built or sold the tool

If you cannot explain how a decision was made without referencing an AI tool, you have an accountability problem.



Turn to your neighbor

90 seconds, be honest!

**What is the boldest claim
an AI vendor has made to
your campus, and what
question do you wish you
had asked?**



A Leadership Framework

Draw the line by consequence, not by complexity





Two tiers for AI use

The tiers are defined by consequence, not technical complexity.

Tier 1: Low-consequence tasks

AI can draft, summarize, suggest, and organize. Errors are easy to catch, cheap to correct, and reviewed before they affect anyone.

Tier 2: High-consequence tasks

Humans decide, review, verify, approve, and remain accountable. Errors have material, reputational, legal, or human consequences.

Same tool. Different governance depending on consequence.

You make the call

For each scenario, vote by show of hands: Tier 1 or Tier 2?

Board memo summarizing financial performance

?

Ask yourself: who acts on this output?

Budget reduction scenarios for cabinet discussion

?

Ask yourself: what happens if it is wrong?

AI screening resumes for a staff position

?

Ask yourself: whose interests are at stake?

The reveal!

Remember: consequence (not complexity) decides the tier.

Board memo summarizing financial performance

Tier 2

Consequential document; every figure and framing choice must be verified.

Budget reduction scenarios for cabinet discussion

Tier 1

Brainstorming before analysis and deliberation.

AI screening resumes for a staff position

Tier 2

People's careers, bias, legal exposure, institutional values.

Judgment is the boundary

AI can help on one side, but humans must lead on the other.

AI can help here

draft
summarize
explore
brainstorm
translate

JUDGMENT

Humans must lead here

decide
approve
verify
take responsibility
explain



Take it back to your team

The framework isn't just for your personal use.

- Where is AI acceptable for preparatory work?
- Where does output require human review?
- Who is accountable for AI-assisted work products?
- What documentation is needed when AI supports a consequential decision?

**Management questions,
not technology questions**



Closing discussion

**What is one decision
on your campus
where AI could assist
but should never
decide?**



Final thought

The technology will keep changing, but the need for leadership does not!

**AI will keep getting more capable.
Your job is to keep hold of the judgment.**

Thank you!



stefani.Langehennig@du.edu



<https://steflangehennig.github.io/>



<https://www.linkedin.com/in/stefani-langehennig-phd-418820144>



Looking to Obtain CPEs?

**Don't forget to
Check-Out from
the system!**

