

From Unstructured Signals to Governed Decision Intelligence

Experiment: Dream Dwell AI Consultant Role Intelligence Workflow

Owner: Manny Sekhon | Founder & Principal Consultant, Dream Dwell LLC

Objective: turn fragmented, unstructured opportunity information into a repeatable decision workflow with explicit criteria, evidence mapping, structured outputs, and human approval.

1. Business problem

Opportunity evaluation is usually fragmented across job descriptions, notes, resume variants, research, and subjective judgment. That creates inconsistent prioritization, repeated manual analysis, weak traceability, and a risk of allowing persuasive AI output to outrun the underlying evidence.

2. Experiment hypothesis

AI can accelerate the analysis, but the operating model must constrain it. The experiment therefore separates intake, normalization, evaluation, evidence mapping, recommendation, and approval rather than asking one model to make an opaque end-to-end decision.

3. Workflow architecture

1	INTAKE	Unstructured opportunity / role information
2	NORMALIZE	Convert inputs into consistent fields and comparable structure
3	AI ANALYSIS	Extract requirements, themes, risks, and decision signals
4	EVIDENCE MAP	Compare requirements against approved profile / resume evidence
5	EVALUATE	Apply defined criteria and structured scoring
6	RECOMMEND	Produce a prioritized decision with gaps and rationale
7	HUMAN GATE	Review, approve, reject, or revise before downstream action

Design principle: AI generates decision support; it does not receive authority to approve its own output.

The experiment is designed as an operating system, not a prompt.

The central question was not simply whether an LLM could produce a useful answer. It was whether the workflow could make AI output more consistent, reviewable, reusable, and safe enough to support real operating decisions.

Control model

CONTROL	PURPOSE
Structured schemas	Reduce output drift and force comparable fields across evaluations.
Defined evaluation criteria	Separate evidence-based fit from persuasive model language.
Evidence mapping	Tie recommendations back to approved source material instead of unsupported inference.
Human-in-the-loop approval	Keep consequential decisions with a person, not the model.
Explicit gaps / uncertainty	Expose missing evidence rather than silently filling it.
Repeatable workflow stages	Make the process testable and easier to improve stage by stage.
Auditability	Preserve the reasoning inputs, criteria, outputs, and approval points needed for review.

What I tested

I tested whether a multi-stage workflow could reliably convert unstructured role information into a structured evaluation while preserving the distinction between **facts, evidence, gaps, and recommendations**. The experiment also tested whether human review could be inserted as an explicit control rather than an informal afterthought.

What changed through iteration

The workflow moved away from generic AI-generated career advice toward deterministic intake, explicit scoring criteria, approved evidence sources, role-specific evaluation, and controlled output. The practical lesson was that better AI outcomes came less from adding more prompting and more from improving the surrounding process, constraints, and review gates.

MODEL CAPABILITY Fast synthesis, pattern detection, drafting, comparison	OPERATING CONTROL Schemas, criteria, evidence, review gates, auditability	HUMAN ACCOUNTABILITY Approve, reject, revise, and own the consequential decision
--	---	--

Why this experiment matters beyond the original use case

The experiment is intentionally domain-transferable. The same pattern can be applied to Customer Success and enterprise transformation wherever teams need to convert noisy signals into governed, prioritized action.

Customer Success translation

EXPERIMENT PATTERN	CUSTOMER SUCCESS / TRANSFORMATION APPLICATION
Unstructured intake	Customer feedback, support signals, usage context, executive goals, implementation issues
Normalization	Consistent customer-health / transformation inputs
AI analysis	Summarize themes, risks, blockers, adoption signals, and opportunities
Evidence mapping	Tie recommendations to observable customer and operating evidence
Evaluation	Prioritize health risks, workflow friction, adoption gaps, and value opportunities
Recommendation	Generate next-best actions, executive discussion points, or intervention options
Human approval	CS / transformation leader validates action before customer-facing execution

Key lesson

The highest-value AI work is often operating-model work. The model matters, but so do the inputs, decision rights, workflow stages, controls, evaluation criteria, feedback loops, and accountability around it. That is the approach I bring to AI transformation: use AI to increase speed and intelligence without sacrificing governance, customer trust, or human ownership.

Current direction

I am continuing to develop this pattern as a reusable decision-intelligence architecture for workflow modernization, executive reporting, risk visibility, prioritization, and AI-assisted operating systems.

Experiment summary

Unstructured signals -> structured evidence -> AI-assisted evaluation -> governed recommendation -> human decision