

DECISION-CENTRED ENTERPRISE UX · INSURANCE

Designing Better Decisions Across First Notice of Loss

How research, systems thinking and responsible decision support improved a complex motor-claims workflow.

My role	End-to-end Product Designer
Focus	Decision quality
Scope	Intake to claim routing

My job wasn't to add AI or produce screens — it was to help the organization make better decisions across claimant intake, evidence review, prioritization and exception handling.

Confidentiality: company name, proprietary rules and exact performance data are omitted. Interfaces shown are reconstructed for portfolio use.

01 · PRODUCT CHALLENGE

The visible form was only one part of the problem

An intake agent once told me she'd typed the same collision details into three different systems in one call — once while the claimant was still speaking, once into a case note, once into the policy system — because none of the tools shared what the others already knew. That repetition wasn't a training gap; it was the shape of the whole workflow.

A motor claim begins when a customer reports a loss. Behind that first interaction, teams must validate policy details, capture the incident, identify injury or severity signals, review evidence, and route the claim correctly.

THE PROBLEM

Fragmented tools and unclear decision ownership slowed intake, created inconsistent records, and pushed avoidable uncertainty downstream.

DESIGN OBJECTIVE

Create one coherent decision journey: capture the right evidence, expose uncertainty early, help each role take the next responsible action, and measure whether those decisions improved claim readiness.

Role	What they needed
Claimant	Clarity and reassurance
Intake agent	Fast, accurate capture
Adjuster	Reliable claim readiness
Supervisor	Queue and exception visibility

KEY CONSTRAINTS

- Regulated decisions and sensitive customer information
- Multiple systems, evidence types, and operating roles
- High emotional load during the first claim interaction
- AI uncertainty that could not be hidden behind automation

02 · DISCOVERY & SYNTHESIS

I investigated where decisions were failing, and why

I worked with claims operations, product, engineering, data and compliance to understand not just how a notice moved, but who made each decision, what evidence they used, and what happened when confidence was low.

- Stakeholder and operations interviews
- Intake and document-review walkthroughs
- Current-state journey and system-dependency mapping
- Form, validation, and exception-message audit
- Prototype-based usability sessions

What emerged	What it meant
Duplicate capture	Incident details repeated across calls, notes and systems
Hidden uncertainty	Users couldn't distinguish source data from AI-inferred data
Late exceptions	Missing or conflicting details surfaced after intake, causing rework
Weak prioritization	Queues didn't consistently expose severity, readiness or service risk
Automation anxiety	Users resisted recommendations they couldn't inspect or correct

Design implication: the experience had to improve decision quality — technology was only one part of the intervention.

03 · OPTIONS CONSIDERED

Guided capture, not full automation

Option	Strength	Risk / limitation	Decision
Rules-based validation only	Simple, predictable, low risk	Missed nuanced injury and liability signals	Not sufficient alone
Fully automated triage & routing	Fastest theoretical throughput	No inspection point before a regulated decision; automation anxiety	Rejected
AI-assisted, human-confirmed intake	Speed with an inspection point at every suggestion	Requires more design work on evidence and confidence UI	Selected

An early instinct was to let the system auto-populate and auto-route straightforward claims to cut agent effort further. Usability sessions changed that: agents didn't resist the suggestions, they resisted suggestions they couldn't check. The fix wasn't less AI — it was making every suggestion inspectable before it ever reached a queue.

04 · DECISION ARCHITECTURE

Separate evidence, interpretation, and accountability

I redesigned the service around decision points rather than system screens. Each stage clarified the evidence available, the person accountable, the support provided, and the consequence of proceeding with uncertainty.

1	2	3	4	5
Report loss (guided intake)	Upload evidence (photos + docs)	AI pre-check (extract + flag)	Human review (confirm + correct)	Route claim (priority + owner)

DECISIONS I MADE

- Reduced early questions to the minimum evidence needed to progress
- Preserved the claimant's narrative beside structured information
- Linked every system suggestion to a visible evidence source
- Moved uncertainty from one opaque score to the affected field
- Designed exceptions as owned states with reasons and recovery paths
- Captured confirmation, correction, and routing rationale for accountability

05 · DECISION SUPPORT

Make every recommendation inspectable

Suggestion	Evidence	Confidence	Control
What AI proposes	Why it proposes it	Where uncertainty exists	Edit, confirm or escalate

WHY AI WAS USED, AND WHERE IT STOPPED

AI was appropriate for structuring evidence, summarizing narratives, and highlighting missing information. It was not appropriate for final coverage, liability, fraud, or settlement decisions. I translated that boundary into visible interaction rules and human checkpoints.

06 · EXCEPTION MANAGEMENT

Turn uncertainty into accountable work

When a policy number differed across two documents, the system didn't silently pick one — it named the conflict in plain language, showed both sources side by side, and required a documented resolution before the claim could proceed.

- Conflicts described in plain language, not error codes
- Competing values retained their evidence source
- Recommendations supported, never replaced, human judgment
- Resolution notes created an auditable record
- Low-confidence output could never silently enter downstream processing

07 · VALIDATION & ITERATION

I evaluated decision quality, not screen preference

What I validated: creating an FNOL from a claimant narrative and documents, finding and correcting an AI-extracted field, explaining why a claim was prioritized, resolving conflicting policy information, and resuming an incomplete notice without losing context.

Before	After	Why it changed
One overall confidence score treated the whole output as equally reliable	Confidence moved to field level with source evidence and explicit confirmation	Users need to know which fields to trust, not just an average
Exceptions appeared as system errors with no owner	Each exception got a reason, owner, recommended action and audit trail	An error state without ownership doesn't get resolved, it gets escalated verbally

08 · WHAT I LEARNED

Trust is a design decision, not a model output

- Automation anxiety usually isn't resistance to AI — it's resistance to suggestions people can't check. Making every value traceable to its source did more for adoption than any accuracy improvement.
- One confidence score hides which specific fields need a second look — moving confidence to the field level is what let agents actually triage their own attention.
- An exception without an owner doesn't get resolved, it gets escalated informally and inconsistently — naming a reason, an owner, and a recovery path turned uncertainty into ordinary, trackable work.
- Reducing repeated capture mattered as much to trust as anything AI did — claimants and agents judged the whole system by whether it remembered what they'd already said.
- Success measured as "volume processed" hid the real problem — reframing it as readiness, accuracy and decision confidence is what got operations, compliance and product aligned on the same goal.

MY 2026 UX CONTRIBUTION

- Framed the business and user problem before selecting a technology response
- Connected user research with operational, risk, and product evidence
- Defined decision ownership, workflow states, and measurable success signals
- Designed the end-to-end service across people, policy, data, and interfaces
- Used AI selectively where it improved judgment and retained human accountability
- Aligned product, engineering, operations, data, and compliance around trade-offs

Design leadership: I reframed success from "how much can AI automate?" to "how consistently can people make the right decision with the right evidence?"

09 · MEASURED IMPACT

A clearer workflow produced stronger operational signals

Figures are rounded and ranged to protect confidential data. They come from a before/after operational comparison across participating intake and claims teams.

Measure	Before	After	Change
Average FNOL handling time	18-22 min	11-14 min	~30-40% faster
Manual field re-entry	6-8 fields	2-3 fields	~55-65% lower
Incomplete notices reaching review	18-22%	9-12%	~40-50% lower
Correct first-time routing	72-78%	86-90%	~10-16 pts higher
Critical conflicts found at intake	45-55%	78-85%	~25-35 pts higher

The value came from better evidence and accountability, not from AI adoption alone.

APPENDIX**How I'd walk through this in the room**

- Context — an agent retyping the same incident details into three systems in one call
- The decision not to fully automate — why guided, inspectable capture beat auto-triage
- The interaction contract — suggestion, evidence, confidence, control
- One trust moment — a policy number conflict resolved in plain language, not an error code
- Outcome — faster handling, less rework, and what I'd do differently next time