

COGPLEXITI · SYSTEMS & DECISION ANALYSIS

THE SECOND SKILL

Why the Qualification System
Trains Recognition and
Assumes the Rest

JAYSON COIL

COGPLEXITI · SYSTEMS & DECISION ANALYSIS

WHITE PAPER

The Second Skill

Why the Qualification System Trains Recognition and Assumes the Rest

JAYSON COIL

Copyright © 2026 by Jayson Coil. All rights reserved.

Published by Cogplexiti LLC · Flagstaff, Arizona · cogplexiti.com

No part of this publication may be reproduced, distributed, or transmitted in any form or by any means, including photocopying, recording, or other electronic or mechanical methods, without the prior written permission of the publisher, except in the case of brief quotations embodied in critical reviews and certain other noncommercial uses permitted by copyright law.

Cogplexiti · Systems & Decision Analysis

ISBN: 979-8-9970824-2-0

First edition · July 2026

Suggested citation: Coil, Jayson. 2026. *The Second Skill: Why the Qualification System Trains Recognition and Assumes the Rest*. Flagstaff, AZ: Cogplexiti.

How to Read This Paper

This white paper belongs to the Cogplexiti Systems & Decision Analysis line and follows the house analytical format: Executive Summary, Support & Analysis, Red Team Assessment, Dissenting Alternative, and Recommended Course of Action. It is the first of four linked analyses; its central competency — meta-recognition as a trainable, certifiable skill — is consumed by *One Patch Is Not One Culture* as the substrate of shared decision doctrine, and stands behind the cognitive-disciplines material in the *Operational Leader's Field Guide*. The entrapment cases in Section 2 are read at the system level, for mechanism, not judgment; the hindsight caution in Section 2.2 governs the whole reading.

Contents

How to Read This Paper	3
Contents	3
1. Executive Summary	4
2. Support and Analysis.....	5
2.1 Recognition and Its Supervisor	5
2.2 The Stretched Spring.....	6
2.3 The Gap Audit	8
2.4 The Correction That Fits the Constraint	10
3. Red Team Assessment	11
4. Dissenting Alternative.....	12
5. Recommended Course of Action	13
References.....	14

1. Executive Summary

The wildland fire qualification system is a recognition pipeline. Position task books certify observed performance of recognized tasks. The S-curriculum builds the inputs to pattern recognition — fuels, weather, topography, fire behavior prediction. Years of position performance build the pattern library itself. The system does this well, and the result is a workforce whose experienced members assess fire ground situations the way naturalistic decision research says experts assess everything: they recognize the situation as typical, retrieve a workable response, and act.

The system assumes a second skill comes along for free: the ability to critique one's own pattern match while it is still running. It does not train this skill, does not name it in doctrine, and has no instrument to measure it. It assumes accumulated experience produces it on its own.

Three decades of research in naturalistic decision making say the assumption is wrong. Cohen, Freeman, and Thompson (1995) demonstrated that this second skill — meta-recognition, the monitoring and correcting of one's own recognitional conclusions — is distinct from recognition, is trainable in hours rather than years, and produces measurable changes in what decision makers notice. Trained officers in their study detected more evidence that conflicted with the working assessment, and the quality of what they noticed went up, not down. The training did not replace expert intuition with analysis. It taught experts to audit the story their intuition had built.

The entrapment record, read at the system level, is partly the bill for this omission. The crews at South Canyon and Yarnell Hill were qualified by every measure the system keeps. Their recognition was certified. What failed in each case was a frame — a coherent, evidence-based, collectively held story about what the fire was doing — that absorbed a growing series of conflicting cues by explaining each one away, until events outran the story. The mechanism has a name in the research literature. The qualification system has no defense against it, because the system never claimed to train the skill that provides one.

The claim of this paper is bounded. It is not that meta-recognition is absent from the workforce; experienced leaders acquire it informally, through particular mentors, near-misses, and reflection. The claim is that the system distributes the skill by chance, cannot say who has it, and has never treated it as what the research shows it to be: a trainable, certifiable competency. The recommended course of action names the skill in doctrine, installs a field-usable drill at defined decision points, and builds the measurement instrument the system currently lacks — piloted and evaluated before it is mandated, because the supporting evidence, while consistent, is thinner than a national requirement should rest on.

2. Support and Analysis

2.1 Recognition and Its Supervisor

The recognitional baseline. The dominant model of expert decision making under time pressure is recognition-primed decision making (Klein 1993, 1998). Experienced decision makers rarely compare options. They recognize a situation as an instance of a familiar type, retrieve the response that type calls for, run a brief mental simulation to check it, and execute. This is not a degraded form of analysis. It is what expertise is: a large library of patterns, indexed by cues, retrieved fast. The fire environment selects for it ruthlessly. A division supervisor who needed comparative option analysis to answer a spot fire across the line would be a hazard.

The meta-level. The Recognition/Metacognition model (Cohen, Freeman, and Thompson 1995; Cohen, Adelman, Tolcott, Bresnick, and Marvin 1993) describes what sits above recognition in proficient decision makers. Recognition operates at the object level: cues activate schemas, schemas assemble into a situation model — a story about what is happening and what will happen next. The meta-level monitors that story. It critiques the story for three classes of problems and corrects the problems it finds. The three classes matter because each fails differently on a fireline:

Incompleteness — the story has gaps. The frame says the fire will remain a low-intensity ground fire, but the story has not accounted for what the forecast wind does to the fuel bed at 1600. A complete story about fire behavior has the same skeleton every time — fuels, weather, and topography producing expected behavior; expected behavior meeting the planned engagement; the engagement producing an outcome — and an incomplete story is one where a component is missing and nobody has noticed (the story-structure logic follows Pennington and Hastie 1993).

Unreliability — the story is complete but rests on assumptions that have not been examined. The escape route is adequate *assuming* the fire holds its current rate of spread. The safety zone is reachable *assuming* the trigger point provides the lead time it was drawn to provide. The story works. The question is what it costs to believe.

Conflict — a cue arrives whose natural, recognitional meaning contradicts the story. Smoke color changes. A lookout reports fire behavior the frame says should not be happening yet. Conflict is the most valuable signal the environment sends and the easiest to spend cheaply, because a coherent frame can almost always buy an explanation for one more discordant cue.

The quick test. The meta-level is itself governed. Critiquing takes time and attention, and the R/M model specifies when to engage it: when stakes are high, when time permits, and when the situation is novel — when the pattern match is least trustworthy. When stakes are low or time is gone, recognition runs uncorrected, and should. This regulating judgment — knowing when the situation has earned deliberate critique — is itself part of the skill, and it maps directly onto the inflection-point logic of operational decision doctrine: there are identifiable moments when the cost of a few minutes of structured doubt is low and the value is highest (Coil 2026).

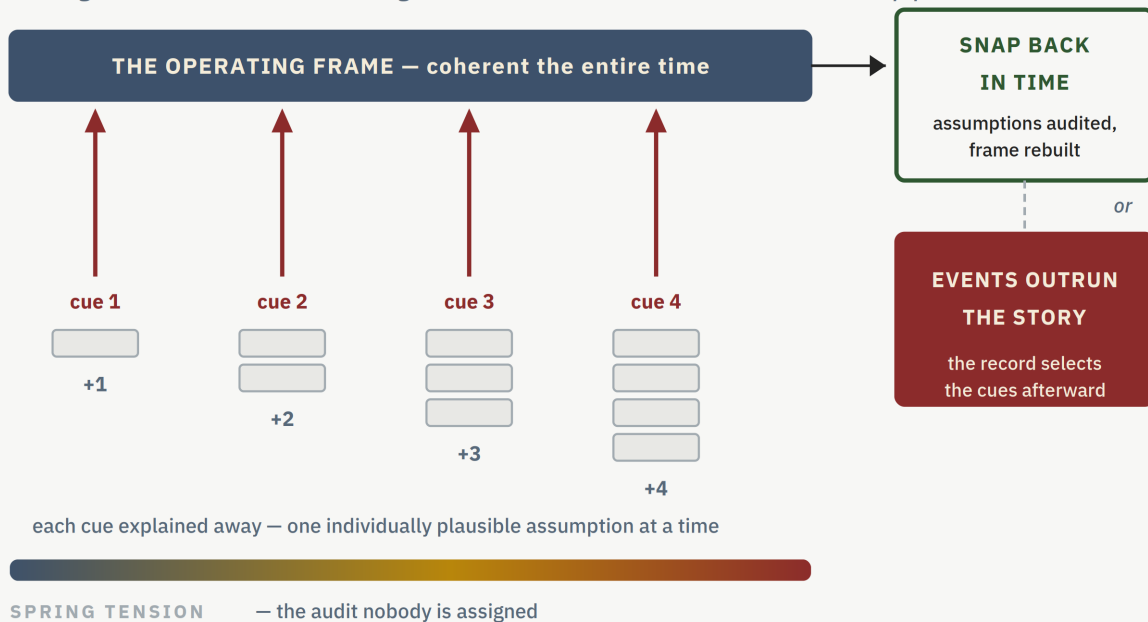
The R/M model's practical claim is modest and precise. Experts do not need to be taught how to recognize; they need a disciplined way to test what recognition hands them, proportioned to time, stakes, and novelty. That is the second skill.

2.2 The Stretched Spring

FIGURE 1

The stretched spring

A coherent frame absorbs conflicting cues by adding assumptions. Every step is locally rational; nothing announces the accumulating tension. The meta-level audit is the skill the pipeline never installs.



The error is not the explaining. The error is losing track of how much explaining the frame now requires.

The mechanism. Cohen, Freeman, and Thompson describe what happens when a decision maker holds a frame while conflicting cues accumulate. Each conflicting cue is explained away with an assumption that is, taken alone, individually plausible. The bridge was destroyed — perhaps the enemy has better bridging equipment than we thought. The artillery moved north — perhaps it will move back south before the attack. Each explanation stretches a spring. The assumptions accumulate, mostly unexamined, and the frame remains coherent the entire time — that is what makes the mechanism dangerous. There is no felt moment of irrationality. Then either the spring snaps back in time — the decision maker registers the accumulated tension, drops the assumptions, and builds a new story — or events resolve the question instead.

Two properties of this mechanism deserve emphasis before it is applied to the fire record. First, explaining away conflicting data is not a bias to be trained out. Cohen (1993) argues, correctly, that abandoning a well-supported frame at the first discordant cue is its own failure mode; a single surprising observation is often wrong, misread, or incomplete, and frame stability is a feature of expertise, not a defect. The error is not the explaining. The error is losing track of how much explaining the frame now requires. Second, the mechanism is silent by design. Every individual step is locally rational. Nothing in the sequence announces itself as a warning, which is precisely why an externally supplied trip-wire list cannot catch it: the tension lives in the decision maker's own assumption stack, and only an audit of that stack reveals it.

South Canyon, read as a spring. On July 6, 1994, fourteen firefighters died on Storm King Mountain when the South Canyon Fire transitioned under a dry cold front and overran crews constructing line below the fire. The investigation record (South Canyon Fire Accident Investigation Team 1994) and the fire behavior analysis that followed (Butler et al. 1998) catalogue the cues now taught in every classroom: drought-cured Gambel oak, a forecast frontal passage with strong winds, line construction downhill toward unburned fuel with fire below, escape routes that ran uphill through green. This paper does not re-litigate those cues, and it does not ask why the people on that mountain failed to see what a classroom can see with the outcome in hand. It asks a system question: what was the operating frame, and what did the system provide for auditing it?

The frame was coherent. The fire had smoldered for days at low intensity. It had been assessed, repeatedly and at multiple levels, as a low-priority incident among many on a stretched district. Within that frame, each cue had an available, individually plausible accommodation — the fire had already had opportunities to make a run and had not; resources were engaged and progress was visible; the forecast was one input among many on a fire that had so far declined to behave like a problem. The frame absorbed the cues one at a time. What the system had supplied to the people holding that frame was recognition training, watch-out lists, and experience. What it had not supplied — to them or to anyone — was a named, trained, practiced method for asking: *how many*

assumptions is this story now carrying, and when did we last examine them? The human factors workshop convened after the fire (Putnam 1995) circled exactly this territory: attention, stress, and the persistence of situation assessments under load. The workshop's findings entered the curriculum as awareness concepts. The audit method never arrived.

Yarnell Hill, read as a spring. On June 30, 2013, nineteen members of the Granite Mountain Interagency Hotshot Crew died when thunderstorm outflow reversed the Yarnell Hill Fire and overran them as they moved through unburned chaparral toward a ranch identified as a safety zone. The Serious Accident Investigation Report (Arizona State Forestry Division 2013) found the crew's actions consistent with standard practice and declined to assign individual fault. That finding is not an obstacle to this analysis. It is the analysis. A fully qualified Type 1 crew — certified recognition, current training, deep experience in the fuel type — operated within norms while holding a frame whose supporting assumptions were individually reasonable and collectively unforgiving: that the fire's direction and rate gave time for the move; that the route's exposure was bounded; that the broadcast weather update meant for the fire what it had meant on other afternoons. The outcome turned on the arithmetic of spread rate against travel time, and the frame's assumptions about that arithmetic were never — so far as the record can show — surfaced as a stack and counted, because counting the stack is not a thing the system trains anyone to do.

The hindsight discipline. Both readings require a caution the reader should hold throughout. We know which cues mattered because the outcome selected them; the same cues appear on dozens of fires every season where the frame holds and nobody writes a report. These cases are used here for the mechanism, not for judgment — they demonstrate that the spring is real, that it defeats qualified people operating within standards, and that it is therefore a system property, not a character flaw. Weick's reading of Mann Gulch (1993) made the structural point a generation ago: when a frame dissolves under surprise and nothing trained stands behind it, what follows is not a worse decision but the collapse of decision making itself. The difference between a spring that snaps back in time and a frame that collapses is not luck. It is whether the meta-level was ever built.

2.3 The Gap Audit

The claim that the system trains recognition and assumes the rest is auditable. Walk the pipeline.

Task books certify recognition. The qualification system (NWCG PMS 310-1) certifies positions through documented performance of tasks under evaluation. The instrument is sound for what it measures: can this person perform the recognized tasks of the position

in real conditions. Nothing in a position task book requires the candidate to demonstrate critique of an active frame — to surface the assumptions a current plan rests on, to notice accumulated conflict, to rebuild a story under time pressure. The behavior is not listed, so it is not observed, so it is not certified.

The S-curriculum improves recognition's inputs. Fire behavior coursework from basic through advanced (S-190, S-290, S-490) makes the pattern library richer and the cues better understood. This is necessary and insufficient: it improves the object level. A better fire behavior forecast inside an unaudited frame is a more convincing story, not a safer one.

The L-curriculum names the territory and stops at the border. Human factors and leadership courses (L-180 forward) teach situation awareness, communication, and error concepts — the closest the pipeline comes. But the treatment is descriptive. Students learn *that* situation awareness degrades and *that* assumptions kill; they are not drilled in a repeatable method for auditing their own running assessment, and no L-course evaluates frame critique as a proficiency with a pass standard. The concepts arrive; the skill is left to form on its own.

Staff rides teach critique in hindsight. The staff ride is the system's most serious engagement with decision failure, and it operates entirely on the wrong side of the outcome. Participants critique a frame whose ending they know. This builds humility and pattern knowledge; it does not build the real-time skill, because the defining difficulty of real-time critique — doubting a story that is still working — is absent by construction.

The lists monitor the environment, not the decision maker. The Ten Standard Fire Orders, the Eighteen Watch Out Situations, and LCES (Gleason 1991) are externally authored trip-wires, and they have saved lives for seventy years. Their limit is categorical: they point outward. They tell the firefighter what conditions in the world should trigger concern. No list points inward at the assumption stack, because a generic list cannot — the stack is specific to this frame, this fire, this afternoon. The Incident Response Pocket Guide's risk management cycle says *reassess*. It does not say how.

The audit's result: recognition is trained, certified, and refreshed on a schedule. Meta-recognition is mentioned in concept, assumed in practice, and measured nowhere. The system therefore distributes the second skill by chance — through the accident of which mentor a leader drew, which near-miss they metabolized, which crew culture they came up in. Some of the workforce has it in depth. The system cannot say who.

2.4 The Correction That Fits the Constraint

A correction has to survive the fireline to matter. Facilitated workshops do not deploy. The R/M training research offers a drill that does.

The crystal ball. The technique Cohen, Freeman, and Thompson trained in ninety minutes runs in minutes and requires no facilitator, no whiteboard, and no group. Take the assessment you are most confident of. A perfect source tells you it is wrong — the fire will not do what you have planned around. Your task is to explain how that could be. The crystal ball rejects your first explanation and demands another, then another. The officers in the study reliably surprised themselves with the number and plausibility of the exceptions they generated — and every exception was an assumption their confident frame had been carrying silently. The drill does not tell you to abandon the frame. Most exceptions, once surfaced, are handled cheaply: some are implausible and dismissed, some are testable with a question to a lookout or dispatch, some drive a small plan adjustment, some are accepted as named risks. The frame usually survives. The confidence in it is now earned rather than assumed.

The conflict variant. The same move inverts when the surprising cue arrives. The crystal ball now insists the frame is *right* despite the cue, and demands explanations. Generating them is fast and clarifying, because the explanations are the price list: the frame's continued validity now costs exactly these assumptions. Trained, the variant compresses to a single habitual question — *how many things have I explained away today?* — which is the spring's tension made countable. Cohen's group found that proficient decision makers already gravitate toward converting all their doubts into this single currency, the reliability of assumptions. The drill makes the conversion deliberate and the currency visible.

Placement. The drill's cost profile fits the pre-commitment window: en route, at transitions, at trigger points — the moments when minutes are available, stakes are identifiable, and the social and cognitive pressures that harden a frame have not yet fully accumulated (Coil 2026). It is the field-weight sibling of the premortem (Klein 2007), stripped to what one supervisor can run in the cab of a truck. Its training burden resembles a briefing format's: modeled, practiced against scenarios, corrected, and made habitual — which is to say, the system already knows how to install skills shaped like this. It has simply never put this one on the list.

The evidence, stated at its actual weight. The experimental support is a controlled study of thirty-seven Army officers showing that ninety minutes of training increased detection of conflicting evidence and improved the rated quality of what officers considered, at conventional and near-conventional significance levels (Cohen, Freeman, and Thompson 1995). That is a real effect in the exact deficit the fire cases exhibit,

produced by a short intervention — and it is one study, three decades old, in an adjacent domain. The mechanism is corroborated across the naturalistic decision making corpus; the training transfer to fire is plausible and undemonstrated. The recommended course of action is built to respect that distinction.

3. Red Team Assessment

3.1 The evidence base is thinner than the claim. The load-bearing experiment is one 1995 study: twenty-nine trained participants, eight controls, effects at $p = .04$ to $.081$, in battlefield situation assessment rather than fire. A skeptical reader — and this paper will find some at the doctrine level — can fairly say that a national training claim is being hung on a pilot-scale result from another domain. The rebuttal has two parts, and only one is fully satisfying. The mechanism rests on the broader NDM corpus, not the single experiment; the R/M model's description of frame maintenance and assumption accumulation is corroborated by thirty years of converging work. But the specific claim that *training transfers to fire ground performance* is an extrapolation. The paper survives this challenge only if the COA is sized to the evidence — pilot, measure, then scale — and it is drafted that way deliberately. If any future version of this argument drifts toward "mandate it nationally now," this red team finding is the correction.

3.2 The skill may already be present and merely unmeasured. Experienced ICs demonstrably do some of this. Tactical decision games and sand table exercises — themselves descended from the same NDM research lineage — already live in parts of the leadership curriculum, and the best operators run informal versions of the crystal ball without a name for it. The sharpest version of this challenge: the paper's deficit framing overstates absence when the real condition is *unstandardized distribution*. The challenge is conceded, and the concession strengthens the argument rather than weakening it. A skill that some of the workforce holds informally, that correlates with the mentors and near-misses a career happens to contain, and that the system can neither identify nor certify, is precisely the definition of a competency the system has failed to institutionalize. The claim is adjusted accordingly throughout: not *absent*, but *distributed by chance and invisible to the qualification system*.

3.3 The entrapment cases carry irreducible hindsight risk. However carefully Section 2.2 is framed, a hostile or grieving reader can compress it to "the paper says the dead should have doubted more." The structural mitigations are in place — mechanism-level analysis, alignment with the investigations' own refusal to assign individual fault, the explicit hindsight caution, the insistence that the spring defeats qualified people operating within standards. They reduce the risk; they do not eliminate it, because the misreading requires only a sentence taken out of context. This is a publication-decision consideration, not a drafting fix, and it should be weighed with the same care given to

adjacent stakes-forward work in the corpus. The alternative — making the argument without the cases — is examined and rejected below, but the rejection deserves to be a conscious choice rather than a default.

3.4 Any mandated drill can decay into theater. The system has watched briefing formats and checklist disciplines erode into recitation. If the crystal ball becomes a line item to initial, it trains nothing and consumes credibility. The countermeasure is in the COA's measurement tier — evaluate the behavior in simulation, never the paperwork — but the risk is named here because it is the most likely long-run failure mode of a successful adoption.

4. Dissenting Alternative

The gap is load, not skill. A credible competing explanation holds that meta-recognition is adequately present in the experienced workforce and is crowded out at exactly the moments it matters — by task saturation, span-of-control failures, communications load, and fatigue. The entrapment record is consistent with this reading: the supervisory positions in both anchor cases were carrying heavy concurrent load. The aviation literature gives the dissent real teeth. Crew resource management research found that monitoring and cross-checking — trained, certified, proficiency-evaluated skills — still degrade first under workload, in crews far more intensively drilled in metacognitive discipline than any fire module will ever be (Dismukes, Berman, and Loukopoulos 2007). If skills that require spare cognitive capacity are the first casualties of load, then a curriculum answer installs a skill precisely where the conditions for exercising it do not exist. On this reading the correct lever is organizational design, not training: protected deputy positions, span discipline, decision windows deliberately shielded from radio traffic — and above all the lookout, which this alternative reframes in a way the paper should adopt regardless of which explanation dominates. A lookout is not merely an environmental sensor. A lookout is institutionalized meta-recognition — a person structurally exempted from the operating frame's workload, positioned to notice what the frame is explaining away. The system built one meta-level position decades ago and has under-theorized it ever since.

The dissent does not win outright, because the load explanation cannot account for frame failures that form *before* load accumulates — assessments hardened at the office, en route, or at morning briefing, in exactly the low-load windows the drill targets. But it wins a permanent seat: any COA that installs training without pairing it to load management is betting against the strongest data the adjacent domains have produced. The recommendation below treats the two as a required pair, not alternatives.

5. Recommended Course of Action

5.1 Name the skill in doctrine. Meta-recognition enters the leadership curriculum as an explicit, defined competency: the monitoring and correcting of one's own recognitional assessments, with working vocabulary — frame, assumption stack, conflict count — stable enough that AARs and peer feedback can reference it. Doctrine is the cheap tier and the prerequisite one; a skill without a name cannot be taught, evaluated, or discussed in a review without improvisation.

5.2 Install the drill at defined inflection points. The crystal ball and its conflict variant are trained the way briefing formats are trained — modeled, practiced against scenario pressure, corrected to habit — and anchored to the decision points where the quick test says deliberate critique pays: en route during the pre-commitment window, at operational transitions and handoffs, and on arrival at named trigger points. The drill is positioned as licensing aggressive action by bounding it, not as institutionalized hesitancy; the output of a two-minute assumption audit is most often a plan executed with earned confidence and two or three named contingencies that did not exist before.

5.3 Measure the behavior, then scale. A pilot population — candidate ICs and division supervisors in existing simulation-based courses is the natural host — is evaluated on observable critique behaviors: assumptions surfaced per scenario, conflicting cues acknowledged versus explained away, frame revisions under injected surprise. Tactical decision games already provide the venue; what is added is a scoring instrument and a proficiency standard. If the measured effect in fire scenarios resembles the 1995 effect in battlefield scenarios, the competency graduates into task book elements for supervisory positions. If it does not, the system has spent a pilot's budget learning something true. National mandate precedes measurement in neither branch.

5.4 Implementation considerations. Three, each inherited from the analysis. The training tier deploys paired with load-management review or not at all — the dissent's evidence on workload is too strong to ignore, and the lookout role should be doctrinally reframed as the system's structural meta-level in the same revision. The measurement tier evaluates behavior in simulation exclusively; the moment the drill becomes documentation, Section 3.4 has occurred. And the competency should be written as shared decision doctrine from the first draft — vocabulary and drills common across legacy cultures — because a merged service in search of an identity substrate can adopt a common way of thinking under uncertainty without any legacy agency surrendering its own history. That argument is made in full elsewhere; the doctrine drafted here should be built to serve it.

End state. A qualification system that can answer a question it currently cannot: *who in this workforce holds the second skill?* — because it named the skill, trained it, and

measured it, rather than assuming experience would deliver it and the record would stay quiet about the cases where it did not.

References

- Arizona State Forestry Division. 2013. *Yarnell Hill Fire Serious Accident Investigation Report*. Phoenix: Arizona State Forestry Division.
- Butler, Bret W., Roberta A. Bartlette, Larry S. Bradshaw, Jack D. Cohen, Patricia L. Andrews, Ted Putnam, and Richard J. Mangan. 1998. *Fire Behavior Associated with the 1994 South Canyon Fire on Storm King Mountain, Colorado*. Research Paper RMRS-RP-9. Fort Collins, CO: USDA Forest Service, Rocky Mountain Research Station.
- Cohen, Marvin S. 1993. "The Naturalistic Basis of Decision Biases." In *Decision Making in Action: Models and Methods*, edited by Gary A. Klein, Judith Orasanu, Roberta Calderwood, and Caroline E. Zsombok. Norwood, NJ: Ablex.
- Cohen, Marvin S., Leonard Adelman, Martin A. Tolcott, Terry A. Bresnick, and F. Freeman Marvin. 1993. *A Cognitive Framework for Battlefield Commanders' Situation Assessment*. Technical Report 93-1. Arlington, VA: Cognitive Technologies.
- Cohen, Marvin S., Jared T. Freeman, and Bryan B. Thompson. 1995. *Training the Naturalistic Decision Maker*. Arlington, VA: Cognitive Technologies.
- Coil, Jayson. 2026. *Operational Leader's Field Guide*. Flagstaff, AZ: Cogplexiti.
- Dismukes, R. Key, Benjamin A. Berman, and Loukia D. Loukopoulos. 2007. *The Limits of Expertise: Rethinking Pilot Error and the Causes of Airline Accidents*. Aldershot: Ashgate.
- Gleason, Paul. 1991. "LCES — A Key to Safety in the Wildland Fire Environment." *Fire Management Notes* 52 (4).
- Klein, Gary A. 1993. "A Recognition-Primed Decision (RPD) Model of Rapid Decision Making." In *Decision Making in Action: Models and Methods*, edited by Gary A. Klein, Judith Orasanu, Roberta Calderwood, and Caroline E. Zsombok. Norwood, NJ: Ablex.
- Klein, Gary A. 1998. *Sources of Power: How People Make Decisions*. Cambridge, MA: MIT Press.
- Klein, Gary A. 2007. "Performing a Project Premortem." *Harvard Business Review* 85 (9).
- National Wildfire Coordinating Group. *Wildland Fire Qualification System Guide*. PMS 310-1. Boise: NWCG.
- Pennington, Nancy, and Reid Hastie. 1993. "A Theory of Explanation-Based Decision Making." In *Decision Making in Action: Models and Methods*, edited by Gary A. Klein, Judith Orasanu, Roberta Calderwood, and Caroline E. Zsombok. Norwood, NJ: Ablex.
- Putnam, Ted. 1995. *Findings from the Wildland Firefighters Human Factors Workshop*. Missoula, MT: USDA Forest Service, Missoula Technology and Development Center.

- South Canyon Fire Accident Investigation Team. 1994. *Report of the South Canyon Fire Accident Investigation Team*. Washington, DC: US Forest Service and Bureau of Land Management.
- Weick, Karl E. 1993. "The Collapse of Sensemaking in Organizations: The Mann Gulch Disaster." *Administrative Science Quarterly* 38 (4): 628–652.

*Experience builds the pattern library.
Nothing in the pipeline trains the auditing of it.*

The wildland fire qualification system certifies recognition — the expert pattern-matching that task books, coursework, and years of position performance build. It assumes a second skill comes along for free: the ability to critique one's own pattern match while it is still running. It does not train that skill, does not name it in doctrine, and has no instrument to measure it. Three decades of naturalistic decision research say the assumption is wrong — and the entrapment record, read at the system level, is partly the bill.

The second skill is trainable in hours, measurable in behavior, and currently distributed by chance. A qualification system that cannot say who has it is assuming what it should be certifying.

The analysis translates the Recognition/Metacognition model to the fire ground; reads South Canyon and Yarnell Hill through the stretched-spring mechanism — qualified people, operating within standards, holding frames whose assumptions were individually reasonable and collectively unforgiving; audits the qualification pipeline for where frame critique should live and does not; and fields the correction that fits the fireline: a minutes-long assumption drill trained the way briefing formats are trained, placed at the decision points where structured doubt costs least and pays most.

The recommended course of action is sized to the evidence: name the skill in doctrine, install the drill at defined inflection points, and measure the behavior in simulation — pilot first, scale on results.

First of four linked analyses in the Cogplexiti Systems & Decision Analysis line. Its central competency is the substrate consumed by One Patch Is Not One Culture, and stands behind the cognitive-disciplines material in the Operational Leader's Field Guide.

ABOUT THE AUTHOR

Jayson Coil is the founder of Cogplexiti LLC, an Assistant Chief with the Sedona Fire District, and a Type 1 Operations Section Chief on Southwest Incident Management Team 1. A U.S. Army veteran and red team instructor, he works at the intersection of operational command and the study of how organizations decide under uncertainty. He has helped develop the systems agencies use to manage major incidents and has presented for the National Fire Academy, the U.S. Forest Service, the Department of the Interior, FEMA, the Haas School of Business at UC Berkeley, and Los Alamos National Laboratory.

