



Climate Sprints

METHOD NOTE · FOR CAPITAL ALLOCATING INTO CLIMATE

How we score whether anyone has to buy it.

The assessment method behind technical and commercial diligence: three assessments, fifteen dimensions, the evidence standard for each, and the rule for what happens when they disagree.

VERSION 1.3
3 SEPTEMBER 2026 · REV 1.3
REVIEW 1Q2027

CLIMATESPRINTS.COM
PRACTICE@CLIMATESPRINTS.COM
POWERED BY PURPOSE. DRIVEN BY ACTION.

What this document is

A fund manager who reads our site and books a call asks one question before any other: how do you actually score commercial compulsion. This is the answer, written down before we are asked.

It exists for three reasons.

To be argued with. A method that cannot be inspected cannot be trusted, and a practice selling judgment that will not show its working is selling assertion. Every dimension below names what evidence counts, what a given score requires, and what would move it. If you think a threshold is wrong, the document is built so you can say which one.

To make assessments comparable. Two companies assessed six months apart by different bench specialists should produce positions that can be read against each other. Without a fixed scale that does not happen, and portfolio-level judgment becomes impossible.

To constrain us. A scale with named anchors makes it harder to talk ourselves into a comfortable answer. The dimensions where we are most likely to be generous — technical stage integrity, and the trigger behind a named buyer — carry the strictest evidence requirements for exactly that reason.

ON THE EXAMPLES IN THIS DOCUMENT

Every technology example here is illustrative. The method is deliberately technology-agnostic and is applied unchanged to software, hardware, biological and chemical processes, and industrial-scale physical infrastructure. What varies between them is the evidence available at a given stage, not the questions asked. Where a dimension's evidence standard cannot be met because of the technology class rather than the company, Section 10 says how that is recorded.

THE LIMIT OF THE METHOD

This is a structure for organizing evidence and stating a position. It is not a model and it does not produce a prediction. Three well-scored assessments still resolve to a judgment made by people who can be wrong, and Section 15 says what we do about that.

The underwriting position

Almost every climate initiative that failed commercially had technology that worked.

What was misjudged was not feasibility. It was whether capital, regulation and buyer urgency would arrive at the same time. A company that is early does not fail gracefully. It runs out of financial or market runway while its thesis is still correct, and the position is written down anyway.

The diligence question is not whether this works. It is who is compelled to buy it, and when they are compelled. If nobody is compelled, the position is a timing bet and should be sized as one.

That produces three assessments rather than two. Technical diligence answers whether the technology functions. Commercial diligence answers whether a market exists. Neither answers whether *this team* can move *this technology* to a *buyer who is compelled to act*, inside the capital and the timeframe available.

The third assessment is the one most often skipped and the one that most often decides the outcome.

The people equation

Underneath all three sits a variable that is the least quantifiable, the most volatile in both directions, and the most consequential.

A founding team that built a technology is, by construction, an inward-facing organization. It has spent years optimizing for a question that has no customer in it. The transition to a market-driven company is not a change of activity. It is a change of orientation, and it asks for skills and a psychological temperament that the founding period actively selected against.

Some teams make that turn. Some see the bottleneck and address it. Some see it and cannot act on it. Some cannot see it at all — and the ones who cannot see it are frequently the most confident that they have already made the turn.

Cohesion, unaddressed ego, unresolved disputes and unexamined blind spots do not appear in a data room. They appear at a buyer presentation, in a negotiation, or in the week a milestone slips — which is exactly when they are most expensive.

Assessment III has been built to surface these before the inflection point rather than after it, and Section 7 is explicit about how, and about the limits of doing so.

A fourth assessment was added in this revision. Capital readiness — whether the company can absorb and deploy the capital it is asking for, over how many rounds, at what dilution, and against what liquidity horizon — is the question a fund manager asks in conventional terms and the one the first version of this method did not answer. Section 8 covers it.

03

Where the established frameworks stop

We did not invent readiness scoring and this method does not replace what exists. It sits on top of three public frameworks and adds what none of them covers.

Technology Readiness Levels

TRL was developed by Stan Sadin at NASA in 1974 and runs from basic research at TRL 1 to full system demonstration at TRL 9.^[1] It is the most widely used readiness metric in existence and the most widely misused.

Two problems for an allocator. First, it measures only technical maturity — it is silent on whether anyone will buy the result. Second, and more practically, it is self-reported. Because TRL is consequential for grant and funding decisions, there is structural pressure to state the number that qualifies rather than the number that is true. The honest reading is that an unverified TRL claim is a conversation starter rather than a conclusion.^[2] A 2015 MIT study of TRL at forty years found persistent inconsistency in how the scale is applied across organizations.^[3]

What we take from it. The stage vocabulary, and the discipline of distinguishing the stage claimed from the stage evidenced. Our TR-1 dimension is anchored to TRL and scored against third-party evidence, not against the deck.

Commercial Readiness Index

CRI was developed by the Australian Renewable Energy Agency in 2014 specifically because TRL stopped short of the questions that decide deployment. It runs six levels from

hypothetical commercial proposition at CRI 1 to bankable asset class at CRI 6, supported by indicators covering regulatory environment, stakeholder acceptance, technical performance, financial proposition on both cost and revenue, industry supply chain and skills, pathways to market, and company maturity.^{[4][5]}

CRI's design question is *can this be financed*. That is a real and different question from whether it works, and the framework is the closest public analogue to what we do.

Where it stops for our purpose. CRI assesses a technology or an application in a market. It does not assess a specific buyer with a specific budget and a specific reason to act this year, and it does not assess the company's capacity to reach them. It was built to help a national agency allocate public funding across a sector, not to underwrite one position in one portfolio.

Adoption Readiness Levels

ARL was developed by the US Department of Energy's Office of Technology Transitions with industry partners, and is operationalized through the Commercial Adoption Readiness Assessment Tool. It assesses seventeen dimensions of adoption risk across four areas — value proposition, market acceptance, resource maturity, and license to operate — grading each low, medium or high, and resolving to a 1-9 score.^{[6][7]}

ARL is the most comprehensive of the three and the most useful to us. Where TRL asks whether the technology works, ARL asks what else has to be true for it to be adopted. It is now applied across DOE demonstration and deployment programs and has been used in published technology assessments.^{[8][9]}

Where it stops for our purpose. ARL is deliberately built to compare across a portfolio of technologies at a category level, and DOE says plainly that it is not intended as a rigorous quantitative analysis of each dimension.^[7] It also assesses the technology's adoption environment rather than the company. A technology class can score well on ARL while the particular company holding it cannot reach a buyer.

WHAT IS GENUINELY MISSING

All three frameworks assess the **technology**. None assesses the **team**, and the venture literature on execution risk is largely anecdotal — practitioner consensus that team quality dominates outcomes, with little agreement on how to assess it and no public scale.^[10]

None of the three assesses capital readiness, and none reconciles its own findings against another framework's. They are run separately, delivered separately, and never argued against each other. That gap is where this method does its work.

What automation has changed

Diligence has automated fast and unevenly. What can be read from documents is now largely machine work. What can only be observed in a room is not — and it is the part with the most variance in it.

We use these tools throughout our own analytical work: market and patent landscapes, financial modeling, document review, pattern analysis across a portfolio. Nothing in this section is an argument against them. It is an argument about what they leave behind.

Where adoption actually is

This is no longer emerging. A 2026 survey of nearly three hundred private capital dealmakers found **85% now using AI** in their workflow.^[17] A separate estimate puts more than 75% of venture reviews as incorporating AI and data analytics.^[18] Practitioner accounts put the data-intensive share of a diligence team's time at **60 to 70%** — and that is the share being compressed.^[19] Automated document and contract review is credited with accelerating that work by 70–80%.^[20]

The phrase everyone uses and nobody specifies

Read across the literature and one sentence recurs almost verbatim. AI augments judgment, it does not replace it. Conviction and relationships stay human. The decision remains human.^{[19][21][22]} The point is put well by one research firm: automation does not reduce the need for expert calls, it increases their value, because the foundational questions are already answered.^[21]

Everyone agrees. Almost nobody specifies what the human part *is*.

"Human judgment" has become the residual category — the name given to whatever has not been automated yet. A residual is not a method. It cannot be assessed, scored, or handed to a colleague.

Three consequences for an allocator

1 • The automated half is converging

When most firms run comparable tools over the same data room, they reach comparable findings. Technical and financial diligence is becoming a commodity input — necessary,

faster, cheaper, and progressively less differentiating. The variance between two funds looking at the same company is migrating almost entirely into the part that was never document work.

2 • Human risk leaves no document trail

This is the structural problem, and it is worse than under-weighting.

A model reads what was written down. The failure modes in Section 8 are precisely those that produce no artifact. A decision deferred past its useful date generates no document — its absence *is* the evidence. A founder who cannot make the shift from inward-facing research to outward-facing commercial work has a calendar that looks busy. A dispute running unresolved for two years appears in no board pack; that is what unresolved means.

So the more diligence is driven by document analysis, the more systematically invisible human risk becomes. Not down-weighted. Absent from the input.

3 • Speed compresses the observation window

Automation shortens timelines, and the pressure is to shorten the whole process rather than to reallocate the time saved. Observation of a team under pressure cannot be compressed the way document review can. It needs sessions, sequence, and a real disagreement to surface. A process that runs in four days instead of three weeks has removed the only conditions under which the third assessment works.

THE POSITION

Technology readiness and market acceptance are now largely assessable by tools, and will become more so. Human capital is the single largest remaining variable, and it is the one the tooling is structurally worst at — not because the models are weak, but because the evidence does not exist in a form they can read.

As diligence becomes more automated, the human component does not become less important. It becomes the only place where judgment still differentiates an outcome, and the only place where a miss is not caught by anybody else's process either.

What this commits us to

Execution capacity is scored from presence, not from documents. Of the twenty-three dimensions, eight sit in Assessment III and none can be completed from a data room. Three of them are scored substantially on how a question is answered rather than on what the answer contains.

Time on it is protected at scoping. Where a compressed timeline makes the session sequence in Section 8 impossible, that is stated as a confidence limit before the engagement

begins rather than discovered at delivery.

No part of it is delegated to a tool. We say publicly that the assessment is machine-assisted and the judgment is not. Concretely: models are used in the analytical work across technical, commercial and capital readiness, and are not used to score execution capacity. The evidence register states which dimensions were machine-assisted on any given engagement.

Where the practice converges

This also answers a fair question about why the practice offers what it offers.

Everything converges on the same variable. Assessment III identifies the human constraint.

Alignment Sprints work directly on it. **Momentum Sprints** address what the constraint is preventing — the commercial motion that has not started because the orientation shift has not happened. The **Climate Catalyst Mastermind** addresses the isolation that is one of its most reliable causes, in a room of people carrying the same weight. Where the gap cannot be closed by support, a **fractional or interim executive** holds the seat until it can.

Section 12 notes that execution capacity is one of only two assessments where a low score is reliably fixable inside a holding period. That is not incidental to how this practice is built. It is the reason.

The line we will not cross

Fractional and interim roles are compatible with independence. They are time-boxed, scoped against a defined outcome, and carry a designed handover. The person in the seat is accountable to the company for delivering something specific. They do not vote and they do not owe a fiduciary duty to the company as a whole.

A board seat is not compatible. A director owes a fiduciary duty to the company. An assessor must retain the ability to write that a team cannot get there, that a technology is a stage behind its claim, or that a position should not be taken. Those two obligations cannot be held by the same party, and holding both would make every prior assessment retrospectively questionable. It also contaminates forward: a practice holding a board seat at one company cannot credibly assess its competitor.

The one exception, stated in advance rather than decided under pressure: a **non-voting observer seat**, taken at the explicit request of the engaging fund, disclosed to all parties, and carrying a standing commitment that we will not assess that company or a direct competitor while the seat is held. Anything beyond that is declined.

Eight design rules

- 1 • Four assessments, never averaged.** Each produces its own score. There is no composite. Section 12 explains why a single number would destroy the finding, and why capital readiness acts as a gate rather than a term in an average.
- 2 • Every score carries an evidence grade.** A 4 supported by a third-party test report and a 4 supported by a founder's assertion are different facts. Section 8 defines the grades and Section 9 the confidence they produce.
- 3 • The absence of evidence is a finding, not a gap.** When a company cannot produce the evidence a dimension requires, that is scored as what it is rather than left blank. A management team that cannot name the budget line has told you something.
- 4 • Named, not estimated.** Dimensions that can be satisfied by a name, a number or a date require one. Addressable market is not evidence of a buyer. A pipeline is not evidence of budget authority.
- 5 • Every dimension states what would move it.** A score with no stated falsifier is an opinion wearing a number. Section 15 collects these into the assessment's own kill criteria.
- 6 • The disagreement is the finding.** When the assessments point different directions, that pattern is the most valuable output of the engagement and it is reported first.
- 7 • The human dimension is never delegated to a tool.** Models are used across the analytical work and are not used to score execution capacity. Section 4 sets out why, and the evidence register states which dimensions were machine-assisted.
- 8 • Every finding becomes an owned entry.** A risk without an owner, a trigger and a review date is a sentence in a report rather than something a fund can act on. Section 13 sets the structure.

Assessment I · Technical readiness

Whether the technology works at the stage claimed, rather than the stage hoped for. Five dimensions, each scored 1 to 5.

DIMENSION	THE QUESTION	EVIDENCE THAT SATISFIES IT
TR-1 Stage integrity	Is the technology at the readiness stage claimed, on evidence rather than assertion?	Third-party test data at the claimed stage. Independent verification of performance under representative rather than ideal conditions. Where a TRL is claimed, the artifact that would substantiate it.
TR-2 Scale-up discontinuity	What physically changes between bench, pilot and first commercial unit, and is any of it unsolved?	Named engineering discontinuities with a stated approach to each. Heat and mass transfer at scale. Materials behavior at duty cycle. Evidence the team has articulated the discontinuity rather than assumed linearity.
TR-3 Dependency exposure	What must be true outside the company's control for this to operate?	Feedstock availability and price exposure. Permitting and interconnection queue position with dates. Single-source components. Where a dependency is critical, evidence of a second path.
TR-4 Performance durability	Does it hold up over the life the cost case assumes?	Degradation curves from operating hours, not models. Uptime under real duty cycle. Maintenance burden. For anything asset-like, evidence over a period proportionate to the claimed life.
TR-5 Independent verification	Who other than the company has looked, and what did they conclude?	Certification against a recognized standard. Customer-run trials with results. Published peer review. Prior technical diligence by a party with no position in the outcome.

Scale anchors

- TR 5 Verified at the claimed stage by a party with no interest in the outcome. Scale-up discontinuities named and closed. Dependencies mapped with alternates. Durability evidenced over a meaningful fraction of claimed life.

- TR 4 Verified at the claimed stage. One or two discontinuities remain open with credible approaches. Dependencies mapped; at least one lacks an alternate.

- TR 3 Works at a stage one step below the claim. Discontinuities identified but not closed. Durability modeled rather than measured.

- TR 2 Functions under controlled conditions. Material discontinuity between current evidence and claimed stage. Key dependency unmapped or single-sourced.

- TR 1 Claim rests substantially on assertion, modelling or analogy to a different system.

A score of 3 or below on TR-1 with an evidence grade of C or D is the single most common finding in early climate diligence and the one most often absent from the deck. It is not disqualifying. It is a repricing.

Assessment II · Commercial compulsion

Not whether a market exists. Whether a specific buyer is compelled to act, and when. This is the assessment the practice is named for.

DIMENSION	THE QUESTION	EVIDENCE THAT SATISFIES IT
CC-1 Named buyer	Is there a person, with a title, at a named organization, who would sign?	Named individual and role. Evidence of contact and of their stated position. Signed LOI, paid pilot or contracted revenue where it exists. A logo is not a buyer; a job title without a name is not a buyer.
CC-2 Budget line	Does the money exist, in whose budget, and in which period?	Identified budget line or capital program. Evidence the buyer has spent against it before. Procurement pathway and approval threshold. Whether purchase is opex or capex, and who signs at that size.
CC-3 Trigger	What forces this buyer to act, and on what date?	A compliance deadline, an asset replacement cycle, a contracted obligation, a supply disruption, a cost threshold already crossed. Dated. A trigger that cannot be dated is a preference, not a trigger.
CC-4 Unsubsidised cost case	Does the economics survive without the support that currently exists?	Cost stack with and without each credit, grant or mandate. Sensitivity to the expiry dates actually legislated. Where the case fails unsubsidised, the volume or input price at which it stops failing.
CC-5 Displacement friction	What does the buyer have to stop doing, and what does that cost them?	Incumbent solution and its remaining book value. Switching cost including requalification, retraining and downtime. Internal advocate and internal opponent. Procurement cycle length measured, not assumed.

Scale anchors

- CC 5 Named buyer, identified budget, dated trigger inside the investment horizon, cost case holds unsubsidised, switching friction quantified and accepted by the buyer.
- CC 4 Named buyer and budget. Trigger dated but outside the near horizon, or cost case holds only at volumes not yet reached.
- CC 3 Buyer identified, budget plausible but unconfirmed. Trigger is directional rather than dated. Cost case depends on support with a legislated expiry.
- CC 2 Interest without budget authority. No dated trigger. Cost case fails without subsidy and the path to parity is not evidenced.
- CC 1 Market described by size rather than by buyer. No named counterparty who has engaged.

THE SINGLE MOST IMPORTANT LINE IN THIS DOCUMENT

CC-3 is the dimension that separates a position from a timing bet. A technology can score 5 on every technical dimension, have a named buyer with a confirmed budget, and still lose money for a decade because nothing forces that buyer to act inside the life of the fund.

When CC-3 cannot be dated, we say so in those words, and we recommend the position be sized as a timing bet. Timing bets are legitimate. They are only dangerous when they are mistaken for something else.

Assessment III · Execution capacity

Whether this team can do what the deck describes. Eight dimensions. The assessment with no public framework behind it, and the one that most often decides the outcome.

We are explicit that this is the least standardized of the four. The venture literature agrees that team quality dominates early-stage outcomes and disagrees on almost everything else about how to assess it.^[10] What structure exists is worth stating: Wasserman's decade-long study of technology and life-science ventures found co-founder conflict implicated in a majority of failures, and found that founder-CEO succession is the norm rather than the exception in high-potential companies.^{[11][12]}

What follows is built from observed failure patterns, structured against that research, and it is the section of this method most likely to change.

The orientation problem

A research organization and a commercial organization are not the same organization running a different play. They select for different people.

Inward-facing work rewards depth, patience with ambiguity, tolerance for repeated failure against a fixed question, and a high threshold for declaring something finished. Outward-facing work rewards speed, comfort with rejection, willingness to ship something imperfect to a buyer, and a low threshold for declaring something good enough to sell.

Those are close to opposite temperaments. A founder who is excellent at the first is not thereby competent at the second, and the belief that the transition is a matter of effort rather than of capability is the most common form the failure takes.

This is not a judgment about the person. It is a fact about the role, and it is addressable — by hire, by fractional support, or by the founder doing genuine work on it. What is not addressable is a founder who will not look.

DIMENSION	THE QUESTION	EVIDENCE THAT SATISFIES IT
EC-1 Demonstrated execution	What has this team actually delivered, at what scale, under what constraint?	Prior shipped products, completed projects, closed rounds, built teams. Verified through reference rather than resume. Operating experience distinguished from advisory experience.
EC-2 Orientation shift	Has the team made, or begun, the move from inward-facing R&D to outward-facing commercial work?	Who in the company owns the buyer relationship, and what have they personally closed. Whether the founder has sat in front of a buyer and been told no. Time allocation of the senior team over the last quarter. Hires or fractional roles made specifically against this shift, and when.
EC-3 Blind-spot permeability	Does the team know what it cannot see, and will it look?	Unprompted articulation of capability gaps. Response when a gap is named by someone else — curiosity or defense. Evidence of a prior belief revised under challenge. A team that names its gaps scores higher than a team with fewer gaps and no awareness of them.
EC-4 Conflict metabolism	How does this team handle disagreement it cannot resolve quickly?	Existence of a stated mechanism — who decides when founders disagree, and on what basis. A traced example of a real disagreement and how it ended. Whether any dispute has been running unresolved for more than two quarters. Equity, title and authority settled or still open.
EC-5 Cohesion and hiring integrity	Is the team hiring for capability alone, or for capability and coherence?	What the last three hires were selected against. Whether cultural and working-style fit was assessed at all. Voluntary departures in the last eighteen months and stated reasons. Whether senior people disagree with the founder in front of others.
EC-6 Decision behavior	How does this team decide when the decision is expensive and the information is incomplete?	Traced examples of prior hard decisions — what was decided, how fast, who dissented, what happened. A decision reversed on new information. A decision deferred past its useful date, and what that cost.
EC-7 Governance permeability	Can a board help early enough to matter, or only after the fact?	Board composition and cadence. Whether bad news reaches the board before it reaches the numbers. Whether prior investors were used as a resource or managed as an audience.
EC-8 Disclosure integrity	Does this team tell you when a number slips, or when you ask?	Prior forecasts compared to prior outcomes. How a missed milestone was communicated and when. Whether the deck's current claims survive contact with the underlying data.

Scale anchors

-
- EC 5** Relevant prior execution at comparable scale. Orientation shift made or credibly underway with named ownership. Gaps articulated unprompted. A real dispute resolved and traceable. Hires selected for coherence as well as capability. Board informed early by habit. Forecast history accurate or accurately explained.
-
- EC 4** Strong execution record with one material gap, named and being addressed. Orientation shift begun. Decision record sound. No pattern of late disclosure.
-
- EC 3** Capable team without directly relevant execution at this stage. Orientation shift acknowledged in language but not yet in time allocation or hiring. Gaps partially recognized. At least one expensive decision visibly deferred.
-
- EC 2** Material capability gap unrecognized, or the orientation shift asserted without evidence. A dispute running unresolved. Bad news reaches the board with the numbers rather than before them.
-
- EC 1** Forecast history contradicts current claims. Founder conflict actively affecting decisions. Defensiveness when a gap is named. Disclosure only under direct questioning.
-

Seeing it before the inflection point

The hardest problem in this assessment is that the failures it looks for are latent. A team can appear coherent through every meeting of a diligence process and come apart in the first hostile negotiation. Interviews conducted in comfortable conditions measure how a team behaves in comfortable conditions.

Three practices, used where the engagement permits:

1 • Stress-tested sessions

At least one session is deliberately run under pressure — a challenged assumption, a compressed decision, a question the team has not rehearsed. The content matters less than what becomes observable: who speaks, who defers, who interrupts, whether disagreement is surfaced or suppressed, and whether the founder can be told they are wrong in front of their team.

2 • Individual before collective

Individual sessions precede group sessions. What people say alone and what they say together is the most reliable signal available on cohesion, and it cannot be obtained in the other order.

3 • The historical trace

Latent conflict is rarely new. Prior disputes, prior departures and prior reversals are traced through reference, with specific attention to how each ended rather than that it did.

THE HONEST LIMIT

None of this is diagnosis, and none of it is reliable at the level of an individual psychological assessment. We report observed behavior under stated conditions and we say what we could not observe.

Where the engagement is post-investment rather than pre-commitment, this work continues under different terms — individual session content confidential, with the fund receiving a position on whether the execution risk has moved rather than a report on any person. That boundary is what makes the work possible.

EC is one of two assessments where a low score is frequently *fixable within a holding period*. A technology that does not work and a market that does not exist are not addressable by portfolio support. A team that has not made the orientation shift, or that has a dispute it has never resolved, sometimes is.

Assessment IV · Capital readiness

Whether the company can absorb and deploy the capital it is asking for — how much, over how many rounds, at what dilution, against what runway, and toward what liquidity event. The conventional question, asked conventionally.

The first three assessments ask whether the thing works, whether anyone has to buy it, and whether this team can get it there. This one asks a different question that a fund manager must answer regardless of the other three: *what does backing this actually commit me to.*

It matters more in this market than it did three years ago. Capital has concentrated — in Q4 2025, 60% of venture capital went to the top ten percent of rounds, and the winners were disproportionately capital-efficient.^[13] Inefficient companies now run out of runway before the next round rather than after it. A company whose plan assumes a funding environment that no longer exists has a capital problem before it has any other kind.

DIMENSION	THE QUESTION	EVIDENCE THAT SATISFIES IT
CR-1 Absorptive capacity	Can this company deploy capital at the rate it proposes without breaking?	Headcount plan against hiring history — actual time-to-fill for the roles proposed. Capital deployed per quarter historically versus planned. Whether the organization has operated at the size the plan reaches. Evidence that spend increases translate into output rather than into overhead.
CR-2 Capital plan integrity	How much capital in total, across how many rounds, and what does each round buy?	Total capital to breakeven, stated. Each round tied to a named milestone that changes the risk profile rather than to a period of time. Bottom-up build of the requirement. Whether non-dilutive sources — grants, debt, customer prepayment, project finance — have been identified and pursued.
CR-3 Runway and milestone alignment	Does the runway reach a milestone that makes the next round fundable by someone else?	Months of runway at current and planned burn. The specific milestone the runway reaches. Whether that milestone is one a later-stage investor recognizes as de-risking. Burn multiple or, pre-revenue, capital consumed per milestone achieved. Target of 18–24 months to the next raise as the working benchmark. [14]
CR-4 Dilution path and cap table	What does the ownership look like at the end of the plan, and does it still work?	Modeled dilution across all planned rounds. Founder and employee ownership at exit. Option pool adequacy for the hires the plan requires. Liquidation preference stack and any structure carried from prior rounds. Whether existing investors can follow on, and whether they intend to.
CR-5 Liquidity horizon	How long to a liquidity event, by what route, and does it fit the fund's clock?	Named route — strategic acquisition, sponsor sale, public listing, project-level monetization. Comparable transactions with dates and multiples. Whether the acquirer set is real and currently acquiring. Time from now to that event against the remaining life of the fund taking the position.

Scale anchors

- CR 5** Absorptive capacity evidenced by prior scaling. Total capital requirement built bottom-up and milestone-tied. Runway of 18–24 months to a milestone a later investor recognizes. Dilution modeled with founders and employees still meaningfully held at exit. Named liquidity route with live comparables inside the fund's horizon.
-
- CR 4** Plan is sound with one soft assumption — usually the round after next, or a liquidity route dependent on a single acquirer. Runway adequate. Cap table clean.
-
- CR 3** Capital requirement is directionally credible but built top-down. Runway reaches a date rather than a milestone. Dilution not modeled beyond the current round. Liquidity route described as a category rather than a set of buyers.
-
- CR 2** Requirement materially understated against comparable companies, or runway under twelve months, or preference stack already exceeding a plausible exit, or no non-dilutive path explored where one plainly exists.
-
- CR 1** Plan assumes a funding environment that no longer exists. No credible total to breakeven. Liquidity horizon beyond any reasonable fund life.
-

WHERE CAPITAL READINESS AND COMMERCIAL COMPULSION COLLIDE

The most dangerous combination in this method is a high CC-3 trigger date that sits *after* the CR-3 runway ends. A buyer who is genuinely compelled — but in thirty months, against eleven months of cash — is not a commercial finding. It is a financing finding, and it converts an attractive position into a bridge decision the fund has not yet priced.

We test this explicitly on every assessment and report it as a dated comparison rather than as two separate scores.

A note on technology class. CR-3's burn multiple benchmark is drawn from recurring-revenue businesses and does not transfer to hardware, industrial or first-of-a-kind infrastructure, where capital consumption per milestone is the more meaningful measure and where the milestone itself is usually a completed unit rather than a revenue level. The dimension is scored on capital consumed per de-risking milestone in those cases, and the substitution is stated in the assessment.

10

Evidence grading

Every dimension carries a letter alongside its score. The letter describes where the evidence came from, not how good the news was.

GRADE	SOURCE	WHAT IT MEANS FOR THE SCORE
A	Independent third party with no position in the outcome. Certification body, customer-run trial, audited data, prior diligence by an unrelated party.	Score stands as written.
B	Primary evidence from a counterparty. Signed LOI, contracted revenue, buyer interview conducted by us, permit granted.	Score stands as written.
C	Company-generated but verifiable. Internal test data with method disclosed, financial model with assumptions exposed, named references we contacted.	Score is reported with the qualification stated.
D	Company assertion, unverifiable within the engagement. Founder statement, undocumented claim, forward projection with no basis supplied.	Score is capped at 3 regardless of the claim.

The cap on grade D is a hard rule. A dimension supported only by assertion cannot score above 3, however plausible the assertion. This is the single mechanism in the method that most reliably prevents a confident founder from producing a confident assessment.

Where evidence could not be obtained

Two situations, reported differently.

Access refused or restricted. Recorded as such, with what was requested and what was withheld. A company that will not permit a buyer interview has produced a finding about CC-1 whether or not it intended to.

Evidence does not yet exist. Recorded as a stage fact rather than a concealment. A pre-pilot technology has no durability data because it cannot; that is scored against TR-4 without any inference about the team.

11

Confidence

Each of the three assessments carries a confidence level derived from its evidence grades, stated separately from the score.

High	No dimension below grade B. Access was complete. Independent verification available on at least two dimensions.
Moderate	Majority at B or above, at most one at D. Access substantially complete.
Low	Two or more dimensions at D, or access materially restricted, or the technology is at a stage where the evidence the method requires cannot yet exist.

A low-confidence assessment is still worth having and is delivered as such. What it is not is a basis for treating the score as settled. Where confidence is low we say what access or evidence would raise it, and what that would cost.

Reconciliation

The four scores are not averaged, summed or weighted into a single number. The pattern between them is the finding.

Why there is no composite score

A technology scoring TR 5 with CC 1 and a technology scoring TR 3 with CC 3 would produce a similar average. They are opposite investment situations. The first is a working technology nobody has to buy — the canonical climate write-down. The second is an ordinary early-stage risk profile with work to do on both sides.

Averaging destroys exactly the information the engagement exists to produce. So we report four scores, four confidence levels, and a stated position on the pattern.

Capital readiness is treated as a **gate rather than a score to be traded off**. A CR of 2 or below does not average against a strong CC. It changes what the position is: no longer an investment in a company but a commitment to a financing sequence, which is a different decision requiring a different reserve.

The divergence patterns

PATTERN	READS AS	POSITION WE TAKE
High TR Low CC	Timing bet. It works. Nobody has to buy it yet.	The most common failure mode in this sector and the one most often mispriced. We say plainly that it is a timing bet, identify what would create compulsion and roughly when, and recommend the position be sized against a longer clock than a standard fund life.
Low TR High CC	Engineering risk with a buyer waiting.	Often the better risk of the two, and frequently under-valued because the deck looks weaker. We assess whether the technical gap is closable with the capital available and state what closing it requires. Compulsion that already exists does not have to be created.
High TR High CC Low EC	Addressable. The thing works, somebody has to buy it, this team is not getting it there.	The only pattern where the deficiency is reliably fixable inside a holding period. We state what specifically must change, whether it is addressable by support or requires a seat filled, and what the cost of each is against the size of the position.
High on all three	Ask why you are seeing it.	Rare, and rarely mispriced. Where we score a company high across all three we say so and then ask the uncomfortable question: what explains this being available to you at this price. Sometimes the answer is good. It should be established rather than assumed.
Low on all three	Pass.	Stated as a pass with the two or three findings that determine it, delivered in a form that can be handed to the company if you choose to.
Strong TR, CC, EC Low CR	A financing decision wearing a company.	The thing works, a buyer is compelled, the team can deliver — and the capital plan does not reach the milestone that makes the next round somebody else's problem. We state the gap in months and dollars, identify the non-dilutive paths not yet pursued, and recommend the reserve be set before the position is taken rather than after.
Compulsion after the runway ends	Right buyer, wrong clock.	CC-3 dates the trigger beyond the point CR-3 says the cash runs out. Reported as a single dated comparison rather than two scores. This is the most frequently missed finding in climate diligence and the one that converts a good company into a down round.
Any pattern, low confidence	Not yet an answer.	We do not convert absent evidence into a position. We state what is missing, what it would take to get it, and whether the timeline of the round permits that.

When our own assessments disagree

Where the technical specialist and the commercial assessment reach findings that cannot both be true, the disagreement is reported as the primary finding rather than resolved internally into a consensus.

The most frequent instance: a technical assessment concluding the system performs at a given cost point, and a buyer interview establishing that the buyer's procurement threshold sits below it. Both may be correctly evidenced. The gap between them is the entire investment question, and averaging it away would be the worst thing we could do with it.

Risk register and mitigation

The assessment produces a position. The risk register turns that position into something the fund can hold, monitor and act on after the wire.

Scores describe a state. They do not tell an investment committee what to watch, who owns it, or what to do when it moves. Every assessment therefore closes with a register built to the structure of the two frameworks a fund's own LPs and auditors already recognize.

What it is built on

ISO 31000:2018, *Risk management — Guidelines*, issued by the International Organization for Standardization, revised from the 2009 edition. It bases management of risk on principles, a framework and a process, and defines risk treatment as the selection and implementation of options to modify risk — avoidance, reduction, sharing, and informed retention.^[15] We use its treatment vocabulary directly.

COSO ERM (2017), *Enterprise Risk Management — Integrating with Strategy and Performance*, issued by the Committee of Sponsoring Organizations of the Treadway Commission. The 2017 revision reoriented enterprise risk management around strategy and performance rather than internal control, and is built on five components spanning governance and culture, strategy and objective-setting, performance, review and revision, and information and reporting.^[16] We use its insistence that risk be tied to a stated objective rather than catalogued in isolation.

Neither framework is certifiable and neither was written for venture positions. What they supply is a common structure and a shared vocabulary, so that a finding from us can enter a fund's existing risk reporting without translation.

How a risk is registered

Every dimension scoring 3 or below generates a register entry. Entries scoring 4 or 5 are registered only where a specific dependency makes them fragile.

FIELD	WHAT IT RECORDS
Source	The dimension that generated it. Every risk traces to a scored dimension; nothing appears in the register that was not assessed.
Objective at risk	What this threatens, stated in the fund's terms — the milestone, the round, the exit, the reserve.
Likelihood	Low, Moderate, High. Assessed against evidence, with the evidence grade carried through from Section 9.
Impact	Contained, Material, or Position-threatening. Material means it changes the return. Position-threatening means it can zero the position.
Treatment	One of four, per ISO 31000: avoid, reduce, share, retain.
Action	The specific thing that would treat it, with an estimated cost and a party who would do it.
Owner	Company, fund, or Climate Sprints. Named. A risk with no owner is not treated.
Trigger	The observable event that says this risk has materialized or receded. Dated where it can be.
Review	When it is next looked at. Tied to a board cycle or a milestone, not to a calendar quarter.

Treatment classes, and what each means here

Avoid	Do not take the position, or take it only with the exposure removed — a carve-out, a milestone-gated tranche, a condition precedent.
Reduce	Act to lower likelihood or impact. Portfolio support, a fractional operator, an independent verification commissioned before close, a hire made a condition of the round.
Share	Move part of the exposure. Syndication, co-investment, insurance where a market exists, contractual allocation to a supplier or offtaker, project-level rather than corporate-level financing.
Retain	Accept it, knowingly, with the reserve sized for it. The legitimate answer for most risks in an early position — but only when it has been named, priced and written down.

THE POINT OF THE REGISTER

Most diligence risk sections are a list of things that could go wrong, which is another way of saying they are unusable. A list has no owner, no trigger and no next review.

The test we hold ourselves to: an entry is well-formed if the fund can put it into a board pack unedited, and a year later can say without argument whether it materialized.

The heat map, and its limits

Entries are plotted on a likelihood-by-impact grid because investment committees read them quickly and it aids triage. Two cautions travel with it.

A heat map compresses. Two risks in the same cell can require entirely different treatments and the grid will not show that. It is a navigation aid to the register, never a substitute for reading it.

And the highest-impact risks in early climate positions are frequently the ones with the least evidence behind their likelihood — precisely the ones a color grid renders most confidently. Where likelihood rests on a grade C or D, the entry is marked as such on the map itself.

Post-investment monitoring

The register is built to be reused. Every entry carries a trigger and a review point, so the same document supports the investment decision and then the holding period.

Where the fund engages us after the wire, the register becomes the working agenda: which entries have moved, which triggers have fired, what has been treated and what has not. Positions fail slowly and visibly far more often than they fail suddenly. A register that is never reopened is a document that predicted the failure and did nothing about it.

14

What you receive

The assessment is delivered as a written position, structured in the same order every time so that assessments can be read against one another.

- 1 **The position.** One page. Which pattern from Section 12, and what we recommend. Written so it can be read aloud in an investment committee.
- 2 **The four scores.** Twenty-three dimensions with score, evidence grade, and the single sentence that supports each.
- 3 **The reconciliation.** Where the assessments agree, where they diverge, and what the divergence means.
- 4 **The risk register.** Every dimension scoring 3 or below, structured per Section 13, with likelihood, impact, treatment, owner, trigger and review. Formatted to go into a board pack unedited.
- 5 **The evidence register.** What we saw, who we spoke to, what was requested and not provided.
- 6 **What would change our view.** Section 14. Written as testable statements, not caveats.
- 7 **If you proceed.** The two or three things to put in the term sheet, the board pack or the first hundred days, drawn from the lowest-scoring dimensions.

Delivered three weeks from the close of the definition phase, followed by ninety minutes with your investment team in which we defend the position and take the argument.

15

What would change our view

Every assessment ends with its own kill criteria: the specific, observable things that would move a score materially, written before the outcome is known.

This is not hedging. It is the mechanism that makes the assessment useful after delivery. A position with stated falsifiers can be monitored; a position without them can only be believed or disbelieved.

Typical form

- A named third party verifying TR-1 at the claimed stage would move it from 3 to 4 and raise confidence from Moderate to High.
- A signed LOI from the named buyer with a delivery date inside eighteen months would move CC-1 and CC-3 together, and change the pattern from timing bet to addressable.
- Legislated extension of the credit the cost case depends on would move CC-4 from 2 to 4. Its expiry on the current schedule would hold CC-4 at 2 and reduce CC-3.
- A commercial hire against the orientation gap named in EC-2, in seat by a stated date, with the buyer relationship transferred to them, would move EC-1 and EC-2.
- Resolution of the founder dispute named in EC-4, evidenced by a written decision mechanism and one decision taken under it, would move EC-4 and raise EC-6.
- A non-dilutive facility of stated size closing before the next equity round would move CR-2 and CR-3 together and reduce the reserve implied by the register.
- Slippage of the CR-3 milestone by one quarter without a corresponding extension of runway would move CR-3 from 4 to 2 and change the pattern to a financing decision.

We also state, once, what would change our view of the *method*. If assessments scoring CC 4 or 5 systematically fail to reach contracted revenue inside the stated horizon, CC-3's dating requirement is too loose and this document is wrong.

16

What this is not

Not an audit. We do not certify financial statements, verify title, or opine on legal or tax matters.

Not a valuation. We take no position on price or terms. Several dimensions bear directly on how a position should be sized, and we say so, but the number is yours.

Not a prediction. The scores describe evidence available at a point in time. They do not forecast an outcome and no score should be read as a probability.

Not investment advice. Climate Sprints is not a registered investment adviser or broker-dealer. Nothing in an assessment is a recommendation to buy, sell or hold any security. We take no equity, no success fee and no compensation from any company we assess.

Not transferable. An assessment is prepared for the client who commissioned it and may not be relied upon by another party without a separate reliance letter.

Not an assessment of the transaction. We assess the company. We do not assess who introduced the opportunity, what their position in it is, or how much of a round reaches the balance sheet rather than retiring an existing holder. Those are material facts about what is being bought, they are frequently invisible in a data room, and they sit with the client to establish. Where we become aware of them in the course of an engagement we report them, but no dimension in this method is pointed at them and you should not assume they have been examined.

Not exhaustive. This is a third-party assessment conducted with the access granted, in the time agreed, on the evidence made available. It cannot see every decision-making vector. Material facts may exist that were not disclosed, not discoverable within the engagement, or not yet in existence at the time of work. The evidence register states what we saw and what we asked for and did not get; it does not and cannot state what nobody knew to look for.

Not a substitute for your own diligence. The assessment is one input to a decision that remains entirely yours. It does not replace legal, technical, financial, environmental, insurance or regulatory diligence conducted by qualified specialists, and it assumes rather than verifies the work of others where that work is relied upon.

No liability for the investment decision. Climate Sprints assumes no liability for any investment decision made by an engaging fund, general partner, family office, developer or other party, or for any outcome arising from a decision informed in whole or in part by an assessment. Liability is limited to the terms of the written engagement agreement.

Time-bound. Every score describes evidence available at a stated date. Regulation moves, competitors ship, key people leave, and credits expire. An assessment more than two quarters old should be treated as a historical document unless it has been refreshed.

Conflicts. We take no equity, no carried interest, no success fee and no referral fee, and we receive no compensation from any company we assess. Where a bench specialist has any relationship with a company under assessment, it is disclosed before work begins and we either substitute the specialist or decline the mandate.

WHO CONDUCTS THE ASSESSMENT

An assessment is not the work of one person. It is led by the founding principal and conducted with the bench specialists the mandate requires — carbon accounting and audit, project development and first-of-a-kind structuring, biochar and waste-to-value, algaculture and biofuels, software architecture and AI governance, ESG data and controls, capital formation and transaction structuring, executive and organizational practice.

The specialist on a mandate is a named participant in the definition phase and in the assessment itself, not a reviewer of someone else's conclusions. Their domain judgment is part of the finding and their disagreement, where it occurs, is reported rather than resolved away.

The specialists engaged on any given assessment are named in the evidence register, with their role and the dimensions they contributed to. Bench specialists are independent professionals engaged per mandate. They are not employees or partners of Climate Sprints, and their participation on one assessment implies no availability for another.

THE RECORD BEHIND THE METHOD

Five decades across climate, cleantech and the circular economy. Eight co-founded companies, two private venture funds, eleven US patents, and a documented record of being right about the technology and wrong about the timing.

That record is not an abstraction. In 1999 the founding principal took a substantial personal position in industrial-scale battery storage — a technology thesis that was correct then and is correct now. He met the principals, saw a demonstration, and read the plan and the projections. He had no instrument for evaluating whether the demonstration evidenced the stage claimed, no way to test whether the projections described a market or reflected his own conviction back at him, and no basis for judging either the founding team or the party who brought him the transaction and whose position his capital largely retired.

The position failed. The larger cost was the reserve it drained and the positions that could not be taken for years afterward.

Each of the three assessments in this method exists because of a specific thing that was not assessable at the time. The full account is published in the founder-facing companion to this document, under his name, because a method built out of a loss is more credible when the loss is stated.

References

Publisher, title, date and URL for every framework and finding cited. Accessed 3 September 2026 · rev 1.3.

-
- [1] Australian Renewable Energy Agency. *Technology and Commercial Readiness Tools* — TRL developed by Stan Sadin at NASA, 1974; scale runs TRL 1 to TRL 9. 10 Jul 2020.
<https://arena.gov.au/knowledge-bank/technology-and-commercial-readiness-tools/>
-
- [2] Innovation Waypoints. *Readiness Metrics: TRL, ARL, CRI, and all the others* — on TRL's susceptibility to self-reporting pressure. 18 Feb 2026.
<https://innovationwaypoints.substack.com/p/readiness-metrics-trl-arl-cri-and>
-
- [3] Olechowski, A., Eppinger, S. D. and Joglekar, N. *Technology Readiness Levels at 40: a study of state-of-the-art use, challenges, and opportunities*. MIT Sloan School Working Paper 5127-15, 2015.
-
- [4] Australian Renewable Energy Agency. *Commercial Readiness Index for Renewable Energy Sectors* — six levels, eight indicators. 2014.
<https://arena.gov.au/assets/2014/02/Commercial-Readiness-Index.pdf>
-
- [5] Carbon Trust. *Assessing the Commercial Maturity of Renewable Energy Technologies: Moving Beyond Technology Readiness* — RE-CRI assessment. May 2017.
<https://ctprodstorageaccountp.blob.core.windows.net/prod-drupal-files/documents/resource/public/>
-
- [6] US Department of Energy, Office of Technology Transitions. *Adoption Readiness Levels* — 17 dimensions across four risk areas, resolving to a 1–9 score.
<https://www.energy.gov/arl/adoption-readiness-levels>
-
- [7] Tian, L., Mees, J., Chan, V. and Dean, W. *Commercial Adoption Readiness Assessment Tool (CARAT)*. US Department of Energy, March 2023. Includes DOE's own statement that the tool is not intended as a rigorous quantitative analysis of each dimension.
[https://www.energy.gov/sites/default/files/2023-03/Commercial%20Adoption%20Readiness%20Assessment%20Tool%20\(CARAT\)_030323.pdf](https://www.energy.gov/sites/default/files/2023-03/Commercial%20Adoption%20Readiness%20Assessment%20Tool%20(CARAT)_030323.pdf)
-
- [8] Pacific Northwest National Laboratory. *Adoption Readiness Level Assessment of Redox Flow Batteries* — applied ARL assessment. Oct 2024.
<https://www.pnnl.gov/publications/adoption-readiness-level-assessment-redox-flow-batteries>
-
- [9] *Aligning sociotechnical systems perspectives and adoption readiness levels to accelerate clean energy adoption*. Nature Reviews Clean Technology, Feb 2026.
<https://www.nature.com/articles/s44359-026-00146-5>
-
- [10] GoingVC. *Risk Assessment for Venture Capitalists* — on team and execution risk as practitioner consensus without a standardized scale. 29 Jan 2026.
<https://www.goingvc.com/post/risk-assessment-for-venture-capitalists>
-
- [11] Wasserman, N. *The Founder's Dilemmas: Anticipating and Avoiding the Pitfalls That Can Sink a Startup*. Princeton University Press, 2012. Longitudinal dataset of technology and life-science ventures; co-founder conflict implicated in a majority of failures.
<https://press.princeton.edu/books/paperback/9780691158303/the-founders-dilemmas>
-
- [12] Wasserman, N. *The Founder's Dilemma*. Harvard Business Review, February 2008 — on founder-CEO succession as the norm rather than the exception in high-potential ventures, and the control-versus-value trade-off.
<https://hbr.org/2008/02/the-founders-dilemma>
-

-
- [13] Carta. *State of Private Markets, Q4 2025* — 60% of Q4 2025 venture capital went to the top 10% of rounds, concentrated among capital-efficient companies. Cited via Waveup, *Capital Efficiency in 2026*, 22 Apr 2026.
<https://waveup.com/blog/capital-efficiency/>
-
- [14] Burn multiple, developed by David Sacks, Craft Ventures. Working benchmarks: below 1.0 strong, 1.0–1.5 efficient at growth stage, above 2.0 a warning; 18–24 months runway to the next raise. Forecastr, *Burn multiple and beyond*, 15 Jun 2026.
<https://www.forecastr.co/blog/burn-multiple>
-
- [15] International Organization for Standardization. *ISO 31000:2018, Risk management — Guidelines*. Principles, framework and process; risk treatment as avoidance, reduction, sharing and informed retention. Revised from ISO 31000:2009.
<https://www.iso.org/standard/65694.html>
-
- [16] Committee of Sponsoring Organizations of the Treadway Commission. *Enterprise Risk Management — Integrating with Strategy and Performance*, 2017. Five components, twenty principles; reoriented from the 2004 internal-control framework toward strategy and performance.
<https://www.coso.org/guidance-erm>
-
- [17] B Capital. *After Decades of Focus on Technology Innovation, It's Time for Climate Companies to Ask: "Can We Scale This?"* — ARL applied to venture assessment. Feb 2026.
<https://b.capital/insights/after-decades-of-focus-on-technology-innovation-its-time-for-climate-companies-to-ask-can-we-scale-this/>
-
- [18] Currence. *H1 2026 Climate Tech Investment & Innovation Report* — Series C share of climate VC rising from 16% to 40%; low-carbon data centers from 3% to 34%. Jul 2026.
<https://www.ctvc.co/h126-climate-tech-funding-up-55-to-26bn-thanks-to-data-centers/>
-
- [17] Affinity. *10 AI tools transforming venture capital in 2026* — survey of nearly 300 private capital dealmakers finding 85% now using AI. Published by a vendor with a commercial interest in the finding. 3 Aug 2026.
<https://www.affinity.co/guides/vc-ai-tools>
-
- [18] StratEngine AI. *AI for VC Due Diligence: Risk Analysis Guide* — estimate that over 75% of venture capital reviews incorporate AI and data analytics. 18 Jan 2026.
<https://stratengineai.com/blog/ai-for-vc-due-diligence-risk-analysis-guide>
-
- [19] Axion Lab. *AI Due Diligence in 2026* — on AI handling the data-intensive analysis that consumes 60–70% of a diligence team's time. 24 Mar 2026.
<https://axionlab.ai/blog/ai-due-diligence-2026>
-
- [20] Qubit Capital. *AI Startup Due Diligence Documents & Metrics Checklist Guide* — automated review accelerating contract and data analysis by 70–80%. 23 Jan 2026.
<https://qubit.capital/blog/ai-startup-due-diligence-documents-metrics>
-
- [21] Third Bridge. *PE due diligence with AI: The complete workflow* — on automation increasing rather than reducing the value of expert calls. 2026.
<https://www.thirdbridge.com/en-us/about-us/media/perspectives/ai-due-diligence-private-equity>
-
- [22] Ricci, T. M. *AI for Venture Capital: 2026 Tools, Costs and Playbook* — on conviction, founder relationships and board-level judgment remaining human. Aug 2026.
<https://www.tommasomariaricci.com/blog/ai-for-venture-capital>
-

A note on source quality

References [1], [3], [4], [6], [7], [8], [9], [11], [12], [15] and [16] are primary framework documentation, standards, or peer-reviewed research. [2], [5], [10], [13], [14], [17] and [18] are practitioner or secondary sources. Where a secondary source materially affects a scoring threshold in this method, the underlying primary document has been consulted directly.

Revision history

v1.0 3 Sep 2026. Three assessments, fifteen dimensions.

v1.1 Execution capacity expanded to eight dimensions with the orientation shift, blind-spot permeability, conflict metabolism and cohesion added. Assessment IV, capital readiness, added. Risk register section added on ISO 31000 and COSO ERM. Liability and scope caveats expanded. Bench participation stated.

v1.3 Section 4 added on what automation has changed and why human capital is the largest remaining variable. Eighth design rule added. Independence boundary on board seats stated. Six references added.

v1.3 Origin of the method stated specifically in Section 16, with the full account published under the founding principal's name in the founder-facing companion. Deal provenance added to Section 15 as a limit of company-level assessment.

v1.4 Planned. Internal review and critique from the bench.
