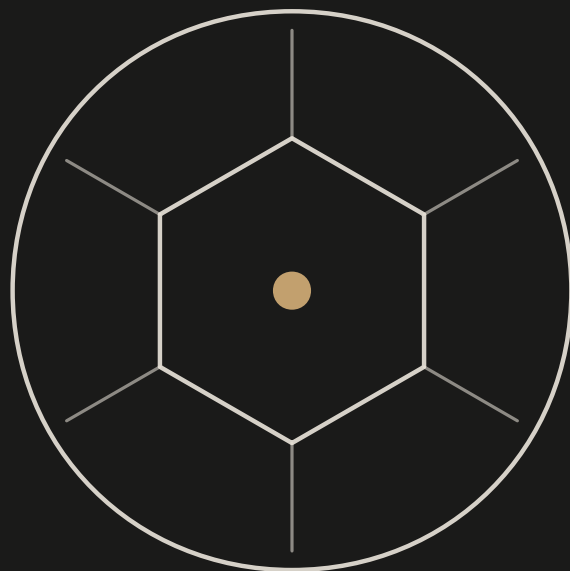




# Aperture

A Technical Manual

The theoretical foundations, curation criteria, and validation  
procedure behind the deck.

John Weng, PhD · Bricolas

METHODOLOGY · 2026

## EXECUTIVE SUMMARY

# Executive Summary

Aperture is a curated deck of 190 wordless images for facilitation, coaching, and group work. This manual documents the method beneath the deck: its theoretical foundations, the criteria by which an image is selected, the projective instrument used to test it, and the validation procedure by which it earns a place.

The method rests on a five-part architecture drawn from object-relations and group-relations theory — projection (Klein), the transitional object (Winnicott), the third thing (Hall), the container (Bion), and the aggregate as a reading of the field (Lawrence's social dreaming) — extended, as Aperture's contribution, from the dyad into group facilitation. Images are judged against five published, sovereign criteria and a sixth deck-level criterion, and each is tested before entry by a three-question projective instrument — ground, activate, elevate — against a pass gate: an image enters only when it reliably produces divergent projections across viewers, including at least one outside the curator's home cultural context. For the go-live deck, 190 images cleared this gate against a pool of 77 respondents.

Aperture's claim is qualitative, not psychometric. It establishes that its images produce projective divergence and that the mechanism operates across cultures; it does not claim norms, population representativeness, or psychometric reliability. The standard is trustworthiness — credibility and transferability (Lincoln & Guba) — and the method names its limits explicitly rather than smoothing them over. The criteria are published, the validation is continuous, and the deck is a living instrument.

---

Aperture is a facilitation instrument, not a clinical or psychometric test. It is distinct from clinical projective testing (the TAT as a diagnostic tool): it furnishes a shared third object for a group to think with, and does not measure, score, or classify its users. It is also distinct from existing facilitation image decks in that its curation criteria and validation procedure are published.

**KEYWORDS**

projective techniques · image-based facilitation · group relations · object-relations theory · Thematic Apperception Test (TAT) lineage · experiential learning · facilitation methodology

**AUTHOR**

John Weng, PhD · Bricolas · ORCID 0000-0002-0415-6980

**VERSION**

1.0 · go-live batch, 2026-05-31

**CITE AS**

Weng, J. (2026). Aperture: A Technical Manual — The theoretical foundations, curation criteria, and validation procedure behind the deck (Version 1.0). Bricolas. <https://doi.org/10.5281/zenodo.20822795>

CONTENTS

# Contents

|    |                                     |    |
|----|-------------------------------------|----|
| 01 | Introduction                        | 5  |
| 02 | Theoretical Foundations             | 7  |
| 03 | The Projective Instrument           | 10 |
| 04 | Curation Methodology                | 12 |
| 05 | Validation Procedure                | 14 |
| 06 | Cultural Portability                | 17 |
| 07 | Limitations and What We Don't Claim | 19 |
| A  | The Instrument                      | 22 |
| B  | Curation Rubric                     | 23 |
| C  | Validation Flow                     | 24 |
| D  | References                          | 25 |

# Introduction

This manual documents the method behind Aperture: its theoretical foundations, the criteria by which images are curated, the projective instrument used to test them, and the validation procedure by which an image earns a place in the deck. It is a reference document, written for the practitioner who wants to understand how the method works and why, and for the academically-minded reader who wants to interrogate it. It is a living document, revised as the deck grows and the research continues.

Aperture is a curated deck of images for facilitation, coaching, and group work. What follows is the method beneath that deck.

Aperture is not a psychological assessment. It does not measure, score, or classify the people who use it. It is a facilitation instrument with genuine theoretical underpinnings, and this manual names them, because naming them is how the method earns a practitioner's trust.

Aperture is also not a mood board. The images are not chosen for beauty, mood, or decorative fit. Each is selected against published criteria and tested, before it enters the deck, for one specific property: that it reliably produces divergent projective responses across different viewers. An image everyone reads the same way is not a good Aperture image, however striking it is. The discipline is in the difference. The chapters that follow are an account of it.

## On the kind of claim this manual makes

A note on scope, stated plainly because it governs everything after it. Aperture's research is qualitative, not quantitative. It studies what an image does in the space between a viewer and their projection: the narrative they construct, where they place themselves, the meaning they attribute. It does not study a trait that can be scored on a scale.

The standards of quantitative psychometrics (construct validity, test-retest reliability, normed scores) are therefore out of scope, not unmet. They belong to a different class of instrument, one that measures the

person. Aperture does not measure the person; it furnishes the third thing a group thinks with. Holding it to a psychometric standard would be a category error.

Where this manual speaks to validity, it means the qualitative kind: whether the method does what it claims and produces genuine projective divergence, whether its criteria are applied consistently across the deck, and whether its findings hold across cultural contexts. In qualitative terms this is the language of trustworthiness rather than reliability (Lincoln & Guba, 1985): credibility and transferability. That is the standard against which the validation procedure in Chapter 5 should be read.

## What this manual documents

The chapters proceed in order: the theoretical foundations the method rests on (Chapter 2), the projective instrument used to test images (Chapter 3), the curation criteria an image must satisfy (Chapter 4), the validation procedure that tests it against respondents (Chapter 5), the method's cultural portability (Chapter 6), and its limitations, the things Aperture deliberately does not claim (Chapter 7). Appendices hold the full instrument, the curation rubric, and the reading list.

# Theoretical Foundations

Aperture rests on a five-part architecture drawn from object-relations and group-relations theory. The five parts are not parallel, not a menu of concepts the method borrows from. They are nested. Each layer presupposes the one before it, and an image doing its work in Aperture is doing all five jobs at once: it receives a projection, holds it, makes it shareable, makes it thinkable, and, in aggregate, reaches toward what a group knows but has not yet said. This chapter takes them in order.

## 2.1 Projection

The mechanism by which an image does its work is projection, not projective identification. The distinction is load-bearing. Projective identification (Klein, 1946) is a two-person mechanism: something is placed into another person, who then feels it and acts on it. An image cannot be acted into. When a viewer deposits onto an image what their inner world needs to put somewhere, the image receives it without being destabilized, because it is not a person and has nothing at stake. This is precisely what makes image work generative rather than defensive: the deposit lands somewhere safe enough to be looked at.

## 2.2 The transitional object

What receives the projection is best understood through Winnicott's (1953, 1971) transitional object, the thing that lives in the potential space between a person's inner world and external reality. The image occupies that intermediate zone: neither fully self nor fully other. Because it sits there, it lets a viewer approach material that would be too threatening if addressed directly. People can speak through an image about what they cannot speak about plainly. The image absorbs what has to go somewhere and gives it back as something the viewer can consider rather than defend against.

## 2.3 The third thing

In dyadic and group work, the image becomes the third thing, the shared object in the middle of the conversation. The term is Donald Hall's (2004), writing about love and shared attention:

most of the time our gazes met and entwined as they looked at a third thing.

A group given a third thing can speak about authority, dependency, conflict, grief, or exclusion through the image rather than directly at one another, which is often the only way such things can be spoken at all.

Hall wrote about two people. The structural move into group facilitation is Aperture's contribution: a literary idea adapted and extended into practice. It carries a corollary the method takes seriously: the quality of the third thing determines the quality of what can be spoken around it. A weak image yields weak conversation. This is why curation (Chapter 4) is not decoration but the load-bearing discipline of the method.

## 2.4 The container

A projection landing is not yet useful. What happens next depends on whether there is a structure that can hold it. Bion's (1961, 1962) container is that structure: here, the facilitator, the frame, the time-boundedness, the rules of engagement. Without containment, projection disperses rather than getting metabolized. The container is what makes the image's contents thinkable rather than overwhelming. Several of Aperture's structural choices are container moves in this precise sense: a single image per round, clear participation cues, a bounded sequence. They exist to hold what the image brings up. These container elements — boundary, authority, role, and task — are the same the group-relations tradition names in the BART frame (Green & Molenkamp, 2005), the lineage in which Aperture's group practice sits.



## 2.5 The aggregate, and social dreaming

The final layer is collective. When a group projects onto a shared image, the aggregate of their projections reaches toward something the field knows but cannot yet articulate, the layer Lawrence (1998) named in his social dreaming work. Aperture takes this seriously as data: the pattern of selections across a session is itself a reading of the field, not merely the sum of individual projections. What surfaces is what the room sees when it looks at itself in the image. This, the aggregate as a reading of the field, is Aperture's second original extension of the inherited theory.

### The architecture as a whole

Stated as a sequence: projection deposits the material; the transitional object holds it; the third thing makes it shareable; the container makes it thinkable; the aggregate reaches toward what the field cannot yet say. Remove any layer and the one above it loses its footing, which is why the method treats the five as an architecture, not an inventory.

# The Projective Instrument

Aperture tests each image, and in live use opens each round, with a fixed three-question instrument. The questions are always asked in the same order, because the order is part of the method. The sequence moves from ground to activate to elevate: it settles the viewer in the literal image, then asks them to place themselves inside it, then asks them to speak from it. Each step reaches signal the previous one does not.

## 3.1 Q1. “What is happening in this image?”

**Ground.** Q1 invites a literal, descriptive reading, and responses are typically convergent across viewers: most people see broadly the same scene. That convergence is the point. Q1 confirms the image is legible and gives the viewer somewhere solid to stand before the projective work begins. It also sets a baseline: the contrast between a convergent Q1 and divergent Q2/Q3 responses is itself a signal of an image's projective depth. Where viewers diverge even at Q1, the image is usually too ambiguous to anchor; where they converge all the way through, it is too literal to project onto.

## 3.2 Q2. “Where might you be in this image?”

**Activate.** Q2 is the strongest projective activator in the set. Asking a viewer to locate themselves inside the image requires projection: there is no literally correct answer, so the answer comes from them. Viewers place themselves in fundamentally different relational positions: observer or participant, central or peripheral, identified with a figure or watching from outside it. This is the question that does the psychodynamic work described in Chapter 2: the relational positioning is where the deposit becomes visible.

### 3.3 Q3. “If this image could speak, what would it say?”

**Elevate.** Q3 personifies the image, and the personification does something the first two questions cannot: it gives the viewer permission to voice declarative statements, about meaning, value, what is true in their world, that they would not make if asked directly. Where Q2 surfaces relational position, Q3 surfaces attributed meaning: values, judgments, organizational metaphor. Speaking as the image, viewers say things they would not say as themselves.

### 3.4 The reserved fourth question

A fourth question exists but sits outside the research instrument: “What does this image know about your organization that hasn't been said aloud?” It is deliberately held back. In a self-administered research setting it carries too much weight and assumes a level of psychological safety the format cannot guarantee. It is retained instead as a signature facilitation framing question for live use, where a facilitator, a container, and a present group can hold what it surfaces. Its exclusion from the instrument is itself a container decision (§2.4): the right question in the wrong container is the wrong question.

### 3.5 Validation of the instrument

The three-question set was validated in April 2026 with two coaching colleagues across three images (a hot-air balloon at sunset; Machu Picchu; a girl on a bicycle at a payphone). The findings held the design: Q1 converged on literal description; Q2 produced genuinely different relational positions; Q3 produced divergent meaning-making that neither Q1 nor Q2 reached. The sequence functioned as a coherent progression, ground to activate to elevate, not three independent prompts. This instrument validation is distinct from the per-image validation procedure in Chapter 5: here the question is whether the instrument works; there, whether a given image does.

# Curation Methodology

This is where “not a mood board” is cashed out. An image does not enter the deck because it is beautiful, moving, or a good fit for a slide. It enters because it satisfies a set of published criteria and then passes the validation procedure of Chapter 5. The criteria are sovereign: they decide, not the curator's private taste. They are published, so the basis of selection is inspectable. A mood board has undisclosed taste behind it; Aperture has disclosed criteria in front of it.

## 4.1 Why photographs: the methodological lineage

The choice of photographic imagery is an inheritance, not a preference. The use of figurative images in projective work descends from the Thematic Apperception Test (Morgan & Murray, 1935; Murray, 1938, 1943) and the eighty years of validated projective research that followed it. Photographs are favored over abstract painting for a methodological reason: group-in-relation-to-its-system work needs enough figural structure to anchor projection onto human situations. A purely abstract image scatters the projection; a figurative one gives it something to land on.

A second lineage runs through experiential learning: McCarthy's 4MAT (1980), after Kolb (1984). The deck takes seriously the role of a strong visual in 4MAT's second quadrant: the right image lets a learner feel the territory of a concept before they have language for it. The image does work that words, arriving too early, would foreclose.

## 4.2 The five sovereign criteria

Every candidate image is judged against five criteria.

1. Projective depth. The image yields multiple, materially different readings across different respondents. This is the one criterion not settled by curator judgment alone: it is what the validation procedure (Chapter 5) directly tests.

2. Affective charge. A felt response that arrives before the mind explains it. Because that response is pre-verbal, it can only be scored as a proxy from a written instrument (see Chapter 7).
3. Relational presence. Figures, postures, gazes, or implied human presence. Not limited to literal human figures: figure-ground relationships, scale asymmetries, substrate identifications, and implied absences all qualify. Only purely decorative images, an isolated flower, an abstract color field, a geometric pattern, fail this criterion outright.
4. Metaphoric register. The image embodies rather than illustrates. An image that illustrates a concept (a ladder for “success”) is too literal to project onto; one that embodies makes the viewer do the meaning-making.
5. Compositional integrity. Frame, light, and depth of field serve the image. Craft matters here not for aesthetics but because a poorly made image breaks the projective spell before it can form.

### 4.3 The sixth criterion: tonal register coverage

A sixth, implicit criterion governs the deck as a whole rather than the individual image: the deck must span a range of tonal and affective registers, not cluster in one. An image can satisfy all five individual criteria and still be redundant if its register is already covered. This is a portfolio criterion, not an image criterion: it is the difference between curating a deck and collecting images.

### 4.4 Scoring

Candidates are scored against the five image-level criteria and sorted into four bands: signature, approved, conditional, reject. Projective depth (Criterion 1) is informed by the validation signal of Chapter 5 rather than scored on inspection alone; the resulting entanglement between the measure and the construct it measures is named explicitly in Chapter 7. The four bands are the published standard. The internal protocol maps them to numeric score ranges, which are not reproduced here.

# Validation Procedure

Every image is tested before it enters the deck, not a sample of images, every image. The procedure asks one question of each: does this image reliably produce divergent projective responses across different viewers? The unit of analysis is the image, not the respondent. This chapter states the procedure as designed, then what was actually run for the go-live batch.

## 5.1 The pass gate (as designed)

An image passes when two conditions are both met: (a) at least one pair of respondents produces genuinely distinct projections, and (b) at least one respondent outside the curator's home (U.S.) cultural context has seen it. An image fails when respondents converge on the same reading across the full respondent set, a signal the image is too literal (a Criterion-4 failure) or too flat (a Criterion-2 failure).

Two responses count as distinct projections when they differ on at least one of: the narrative constructed (Q1), the relational position taken (Q2), or the attributed meaning (Q3). The comparison is made on the gestalt of the three responses together, not on any single question: two viewers may both say “ruins” at Q1 yet place themselves in wholly different positions at Q2, which counts as divergent.

## 5.2 Adaptive sampling (as designed)

Testing is adaptive, to spend respondent attention where it is needed.

- Round 1, three respondents. One or more divergent pairs: the image passes, pending the cultural condition. Zero divergent pairs: it advances to Round 2.
- Round 2, up to five respondents. One or more divergent pairs: it passes. Zero across all five: it fails, and the failure type is logged (convergence or blankness) so it informs future curation.

A two-cluster split (for example, three respondents read the image one way, two another) still counts as at least one divergent pair, and passes.

### 5.3 Respondents and recruitment (as run)

The respondent pool is a convenience / purposive sample, recruited in two venues. This is the appropriate sampling method for the kind of claim Aperture makes (Chapter 1): the goal is to establish that images produce divergence, not to estimate parameters of a population, so statistical representativeness is neither claimed nor required.

- Network venue (volunteer, unpaid): trusted colleagues and professional-network referrals, group-relations practitioners, OD practitioners, coaches, leadership-development professionals.
- Prolific venue (paid): recruited to add breadth and, in particular, deliberate routing toward respondents outside the curator's home cultural context.

The figures below are frozen to the go-live batch, as of 2026-05-31.

- Unique human respondents: 77 (the Respondents table holds 79 records, less 2 AI reference respondents).
- Of these, 20 through the Network venue (volunteer colleagues, including the curator) and 57 through Prolific.
- Images in the launch deck: 190, promoted from 191 candidates that passed research validation. This is past the 100-image go-live threshold the protocol set.

### 5.4 Cross-cultural routing

Cultural-portability validation is not a separate phase; it is built into the per-image pass gate (§5.1b) and enforced by routing. When an image accumulates responses without a non-U.S. respondent among them, it is preferentially surfaced to non-U.S. respondents until that condition is met. This makes cultural validation continuous and structural rather than a discrete second study, a point developed further in Chapter 6.

## 5.5 Divergence coding and the AI reference respondent

Divergence is coded by the curator, reviewing each respondent pair's Q1/Q2/Q3 responses side by side and judging the gestalt as convergent or divergent. Coding is done on a rolling basis as responses arrive.

A language model (Claude) serves as a fixed reference respondent on each image, responding to the same three questions before any human responses are collected, to avoid anchoring. The AI response is not counted toward the human divergence threshold or the non-U.S. condition; it functions as a reference point: when every human converges with the model's reading, that is a stronger convergence signal than human agreement alone.

## 5.6 Curator-as-respondent (disclosure)

During the earliest research, before any international respondent had joined, the curator responded to images using the full instrument, logged as the first human respondent. The curator stepped out of the respondent pool at the entry of the first international respondent, to preserve the independence of any cross-cultural claim. This is disclosed again, with its methodological reasoning, in Chapter 7.



# Cultural Portability

Aperture is used across cultural contexts, and the natural objection is that an image meaningful in one culture may be opaque or loaded in another. The method's answer rests on a theoretical layering and one structural refusal.

## 6.1 The theoretical basis

Portability does not rest on the claim that an image means the same thing everywhere. It rests on the claim that the projective mechanism operates across cultures. This is supported by the cross-cultural projective-testing literature descending from the TAT tradition, and by Kardiner's (1939) work on shared symbolic objects across cultures. What travels is not a meaning but a capacity: people of different backgrounds can each deposit their own material onto the same image.

## 6.2 The structural refusal

The safeguard is a refusal built into the method: Aperture does not interpret what a respondent sees. Meaning travels with the respondent, not with the image. An image cannot carry a culturally specific “correct” reading if the method assigns no reading at all. This connects directly to the mechanism of Chapter 2: the image has nothing at stake, makes no claim, and returns to each viewer what that viewer brought. The image is never the meaning; the meaning is what the respondent brings.

## 6.3 Routing as portability in practice

The refusal is a stance; the routing is what operationalizes it. As described in §5.4, an image that has not yet been seen by a respondent outside the curator's home cultural context is preferentially surfaced to such respondents before it can pass. Portability is therefore tested

per image, continuously, rather than asserted for the deck as a whole or deferred to a separate study.

## 6.4 The bounded claim

The portability claim is bounded by the international composition of the respondent pool, which is growing rather than fixed. Aperture does not claim universal cultural portability; it claims that the projective mechanism is cross-culturally operative and that each image has cleared the non-U.S. condition against an internationalizing pool. The honest limit of this claim is stated in Chapter 7.

## 6.5 Removal policy

Cultural portability is not only a gate at entry; it is a standing commitment. An image identified as culturally insensitive in use is removed under a published removal policy (via [curation@aperturedeck.com](mailto:curation@aperturedeck.com)). This is consistent with the deck being a living instrument: the criteria are published, the validation is continuous, and the removal path is open.

# Limitations and What We Don't Claim

This chapter names the boundaries of the method. It is placed last and given weight deliberately: a method that states what it does not claim is more trustworthy than one that does not, and the discipline of naming the limits is itself part of the method.

## 7.1 The claim is qualitative

As stated in Chapter 1, Aperture's research is qualitative. It establishes that images produce projective divergence and that the projective mechanism operates across cultures. It does not establish, and does not attempt to establish, construct validity, test-retest reliability, or normed scores in the quantitative-psychometric sense. Those belong to instruments that measure the person; Aperture furnishes a third thing a group thinks with. The limit is a category boundary, not a shortfall.

## 7.2 Projective depth: measure-construct entanglement

Criterion 1 (Projective Depth) is partly operationalized through the very divergence signal used to validate images, so the construct being measured is partly defined by its measure. This is named openly rather than smoothed over. Aperture treats it as bootstrap co-validation in the sense of Cronbach & Meehl (1955): the construct and its measure are validated together, with the theoretical architecture of Chapter 2 providing the independent anchor that keeps the loop from being circular. A future methods paper that wished to break the entanglement would need an independent measure of projective depth for triangulation.

## 7.3 Affective charge can only be measured as a proxy

Criterion 2 (Affective Charge) asks for a response that arrives before the mind explains it, a pre-verbal event. A written instrument captures only what has been articulated, which is by definition post-verbal.

Affective charge is therefore scored as a proxy, inferred from the scene and from what the response set suggests about the image's capacity, not measured directly. Any report drawing on this data should carry the limitation.

## 7.4 Relational presence is read broadly

Criterion 3 (Relational Presence) is not confined to literal human figures. Figure-ground relationships, scale asymmetries, substrate identifications, and implied absences all satisfy it. This is a scope clarification more than a limit: the criterion should not be read so narrowly that only images of people qualify. Only purely decorative images fail it.

## 7.5 Divergence coding is curator judgment

Divergence is coded by the curator on the gestalt of each respondent pair's responses. Aperture does not present a formal inter-rater reliability statistic as the warrant for its claim, consistent with the qualitative posture of Chapter 1. A language model serves as a reference respondent (§5.5), not as an independent human second coder. The warrant for the claim is the published criteria, the transparent procedure, and the fact that divergence is directly observable in the response record, not a reliability coefficient.

## 7.6 The curator was an early respondent (disclosure)

Before any international respondent joined the pool, the curator responded to images using the full instrument, logged as the first human respondent. The curator stepped out of the respondent pool at the entry of the first international respondent, to preserve the independence of any cross-cultural claim. Pre-international curator participation is acceptable with disclosure because the pre-international research made no cross-cultural claims. This is disclosed in any report drawing on the data.

## 7.7 The empirical claim is bounded

The cultural-portability claim is bounded by the international composition of the respondent pool, which is growing rather than fixed. Aperture is an early-stage, continuously validated deck, not a closed and normed instrument. What Aperture claims: a theoretically grounded method, published curation criteria, and images that have passed the divergence gate against an internationalizing pool. What it does not claim: norms, population representativeness, universal cultural portability, or psychometric validity and reliability.

## 7.8 What we don't do

A set of standing commitments that bound the practice as much as the method.

- We don't A/B-test the covenant. What is offered is what is offered, not what is being optimized against engagement.
- We don't sell analytics on participants. The curator watches selection patterns to inform what enters and retires from the deck. The deck improves; the data is not the product.
- We don't train AI on saved selections. Selections are facilitator-owned.
- We don't keep undisclosed taste behind the deck. The criteria are sovereign and they are published.

## APPENDIX A

# The Instrument

Administered in fixed order; the sequence (ground to activate to elevate) is part of the method (Chapter 3).

- Q1. What is happening in this image? Ground; narrative; typically convergent; confirms legibility.
- Q2. Where might you be in this image? Activate; relational positioning; strongest projective activator.
- Q3. If this image could speak, what would it say? Elevate; personification; attributed meaning.

Reserved, for live facilitation only, not part of the research instrument:

- Q4. What does this image know about your organization that hasn't been said aloud?

## APPENDIX B

# Curation Rubric

Five image-level criteria, each scored; a sixth governs the deck as a whole. The rubric is specified in the internal research protocol; this is its external statement.

| # | CRITERION                            | PASSES WHEN...                                                                                                              |
|---|--------------------------------------|-----------------------------------------------------------------------------------------------------------------------------|
| 1 | Projective depth                     | the image yields multiple, materially different readings across respondents (tested by the validation procedure, Chapter 5) |
| 2 | Affective charge                     | a felt response arrives before the mind explains it (scored as a proxy, §7.3)                                               |
| 3 | Relational presence                  | figures, postures, gazes, or implied human presence, including figure-ground, scale asymmetry, and implied absence          |
| 4 | Metaphoric register                  | the image embodies rather than illustrates                                                                                  |
| 5 | Compositional integrity              | frame, light, and depth of field serve the image                                                                            |
| · | Tonal register coverage (deck-level) | the image extends, rather than duplicates, the deck's range of registers                                                    |

Scoring bands: signature, approved, conditional, reject. The internal protocol maps these to numeric score ranges; the bands themselves are the published standard.

## APPENDIX C

# Validation Flow

Per-image, adaptive. The pass gate and the two adaptive rounds:

PASS GATE ·  $\geq 1$  DIVERGENT PAIR AND  $\geq 1$  NON-U.S. RESPONDENT

---

## ROUND 01 · 3 RESPONDENTS

- $\geq 1$  divergent pair,  $\geq 1$  non-U.S. → PASS · ENTERS THE DECK
- $\geq 1$  divergent pair, 0 non-U.S. → HOLD · CULTURAL PRIORITY
- 0 divergent pairs → ROUND 02

## ROUND 02 · UP TO 5 RESPONDENTS

- $\geq 1$  divergent pair,  $\geq 1$  non-U.S. → PASS
- $\geq 1$  divergent pair, 0 non-U.S. → HOLD · CULTURAL PRIORITY
- 0 divergent pairs (of 5) → FAIL · CONVERGENCE / BLANKNESS

A “Hold · Cultural Priority” image is preferentially routed to non-U.S. respondents until the cultural condition is met; it does not fail on that account.



## APPENDIX D

# References

- Bion, W. R. (1961). *Experiences in Groups and Other Papers*. London: Tavistock Publications.
- Bion, W. R. (1962). The psycho-analytic study of thinking. II. A theory of thinking. *International Journal of Psychoanalysis*, 43, 306-310.
- Cronbach, L. J., & Meehl, P. E. (1955). Construct validity in psychological tests. *Psychological Bulletin*, 52(4), 281-302. <https://doi.org/10.1037/h0040957>
- Green, Z., & Molenkamp, R. (2005). The BART system of group and organizational analysis: Boundary, authority, role and task (Faculty Scholarship No. 78). University of San Diego, School of Leadership and Education Sciences. <https://digital.sandiego.edu/soles-faculty/78>
- Hall, D. (2004). The third thing. *Poetry* (November 2004).
- Kardiner, A. (1939). *The Individual and His Society*. New York: Columbia University Press.
- Klein, M. (1946). Notes on some schizoid mechanisms. *International Journal of Psychoanalysis*, 27, 99-110.
- Kolb, D. A. (1984). *Experiential Learning: Experience as the Source of Learning and Development*. Englewood Cliffs, NJ: Prentice-Hall.
- Lawrence, W. G. (Ed.). (1998). *Social Dreaming @ Work*. London: Karnac Books.
- Lincoln, Y. S., & Guba, E. G. (1985). *Naturalistic Inquiry*. Beverly Hills, CA: Sage.
- McCarthy, B. (1980). *The 4MAT System: Teaching to Learning Styles with Right/Left Mode Techniques*. Barrington, IL: EXCEL.
- Morgan, C. D., & Murray, H. A. (1935). A method for investigating fantasies: The Thematic Apperception Test. *Archives of Neurology and Psychiatry*, 34, 289-306. <https://doi.org/10.1001/archneurpsyc.1935.02250200049005>
- Murray, H. A. (1938). *Explorations in Personality*. New York: Oxford University Press.
- Murray, H. A. (1943). *Thematic Apperception Test Manual*. Cambridge, MA: Harvard University Press.

- Winnicott, D. W. (1953). Transitional objects and transitional phenomena. *International Journal of Psychoanalysis*, 34, 89-97.
- Winnicott, D. W. (1971). *Playing and Reality*. London: Tavistock Publications.