

What Would Convince You? — Ranking the Evidence That Moves Bio-Material Decisions

Led by

Arta Yazdanseta, RPI

BACKGROUND

People already trust this kind of evidence everywhere else. Before booking an Airbnb or buying almost anything online, we scroll past the listing to the reviews, because how a place actually performed for the people who used it is what decides the choice. For one of the largest and longest-lived decisions anyone makes, what to build a building out of, that review layer does not exist.

The gap is measurable. In a systematic review I co-authored of 309 hempcrete publications from 2004 to 2024 (Yazdanseta & Mi, 2025), material characterization filled 78.9% of the literature while social and occupant dimensions accounted for just 2.92%. Hemp-lime is well proven in the laboratory and almost unstudied in the homes people actually live in. Lacking that evidence, the builders, owners, developers, lenders, and policymakers who decide whether these projects happen fall back on cost and energy proxies that misstate both the risk and the value of building this way.

The Habitat Atlas closes that gap. It is the review layer for biomaterial buildings: it turns long-term occupant experience into a queryable decision tool that shows each stakeholder the evidence relevant to their own decisions, with a public MVP planned for summer 2027. An eight-household European pilot has already shown the method produces decision-grade findings, including unanimous re-choice and measurable respiratory improvement among occupants.

The Summit was a chance to pre-validate the Atlas decision layer with the exact professionals the tool is meant to serve, before the platform is built. I used it to answer two questions: whether the framework reads clearly to a working professional seeing it cold, and how different decision-makers weight the same building. Both feed directly into the next iteration of the platform and the next round of occupant interviews.

WHAT I DID

I ran a self-guided station built around a wall-mounted Atlas, which set out the full framework of eight categories and twenty-five subcategories, and two short paper activities.

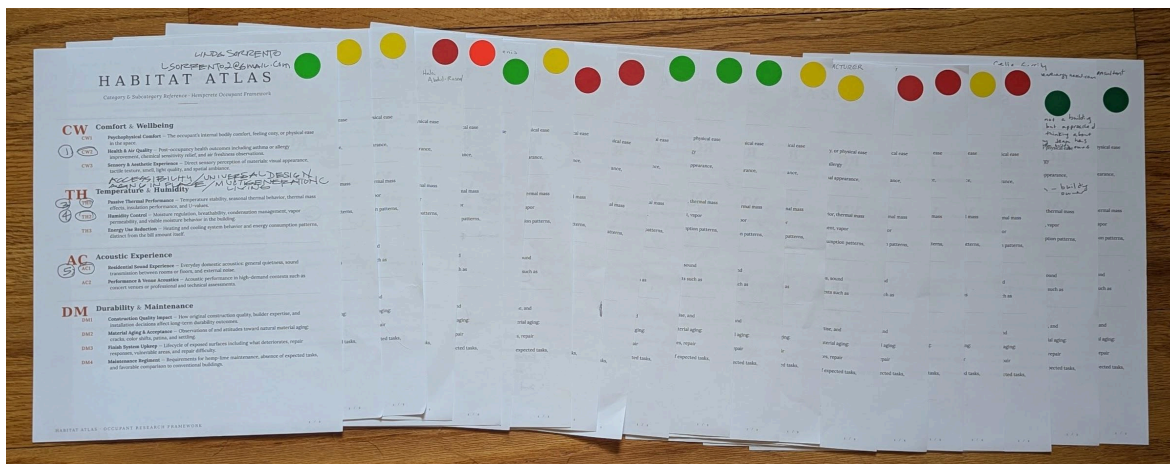


BBMC

Authored by

MASS.

2026 Summit Proceedings



First Activity - Selecting what matters to each stakeholder

The first activity moved through three steps. Working alone, each participant identified their role with a colored dot and circled the five subcategories that weigh most on their own decisions. They then talked those choices over at a table of mixed roles, guided by a prompt asking what the framework was missing for someone in their seat: a category, a question worth putting to future occupants, or a filter they would apply to the data. That discussion did two jobs. It gave people time to test their own picks against other perspectives, which tends to sharpen judgment and surface a more considered answer than a silent form would on its own. And it primed the final step, where each participant wrote down what they felt was missing, along with any other feedback they wanted to leave. The circled selections produced the count data on the Stakeholder Priorities Diagram; the written notes produced the qualitative comments analyzed alongside it.



BBC

Authorized by

MASS.

2026 Summit Proceedings



WHAT THE ROOM SAID

Three results matter most for the next version of the tool. The people who came converge on health and thermal performance, they diverge almost everywhere else, and the priorities that actually drive their decisions are not the ones a simple headcount would surface. A fourth result is structural, and it is about who did not show up.

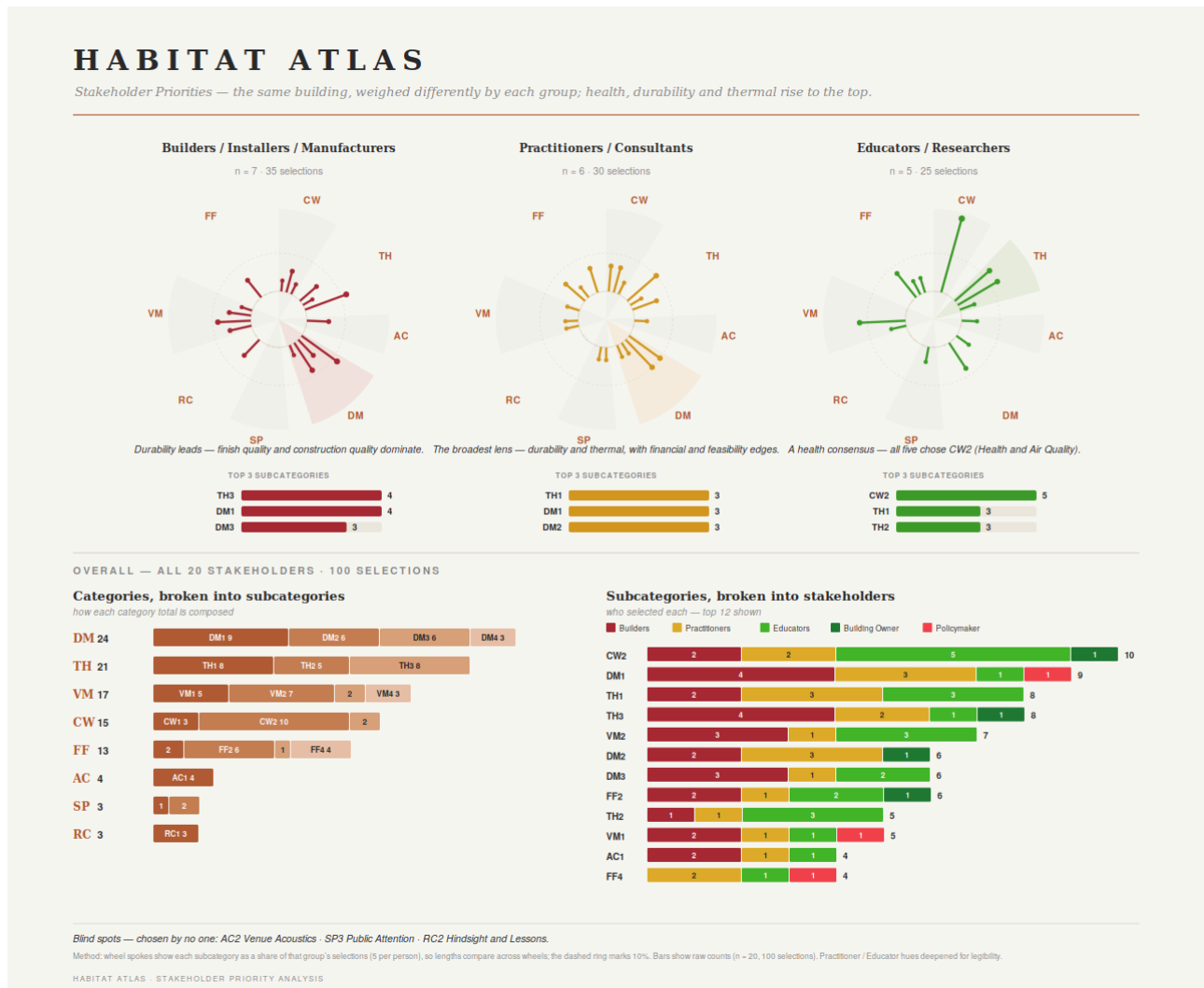
Thermal performance and health are the shared ground. Temperature & Humidity is the top category on the Rank-Weighted sheet (Categories by priority: TH 53, ahead of DM 51 and VM 49) and the



BBMC Authored by **MASS.**

2026 Summit Proceedings

second-highest by count on the Stakeholder Priorities Diagram (Categories, broken into subcategories: DM 24, TH 21). Health & Air Quality is the single most-selected subcategory by count (Subcategories, broken into stakeholders: CW2, 10 selections) and the highest-scoring by rank-weight (Subcategories by stakeholder: CW2, 27). Every educator and researcher chose it, and it tops the builder group on the rank-weighted sheet. For educators, humidity sits right alongside health at the top, so the convergence reads as a health-and-thermal cluster that holds even across groups that disagree on much else.



First Activity - results and findings

The same building looks different through each lens, which is the premise the Atlas is built on, and the radar profiles at the top of the Rank-Weighted sheet show it directly. Builders spread their attention widest and lean on durability, craft quality, and ecological values. Practitioners give the flattest profile of any group, with thermal comfort sharing the front with cost and feasibility. Educators cluster tightly on health and humidity. The two building owners cared most about what the building costs to run and

maintain, and the single policymaker led with ecological ethics, though both of those readings rest on one or two voices.

Priority is not the same as frequency, and the cost finding is where that matters most. Cost Savings is middling when you count picks (FF2, 6 selections, under Subcategories, broken into stakeholders) but jumps to second overall when you weight by intensity (Subcategories by stakeholder: FF2, 26), because a small number of people, weighted toward the building owners, placed it at the very top of their wheel. Energy Use Reduction runs the opposite way: many selected it (TH3, 8 selections, same panel) but it carried far less weight than its frequency suggested (17, mid-pack on Subcategories by stakeholder), a welcome bonus rather than a deciding factor. A tool that ranked priorities by how often they were chosen would get both of these wrong.

The framework itself held up. Participants placed themselves in the taxonomy without coaching, and what they asked for were additions, not repairs. Three subcategories drew nothing from anyone in either activity: Performance & Venue Acoustics, Public Attention, and Hindsight & Lessons Learned, which the Rank-Weighted sheet lists under "Chosen by no one." The whole social-perception and reflections wedge of the wheel stayed nearly bare. The written comments, on the other hand, named gaps that fit the Atlas method and recurred across participants: what these buildings cost to live in over time, how disruptive the build and move-in are, and what happens at the end of a building's life, with accessibility and aging in place raised once.



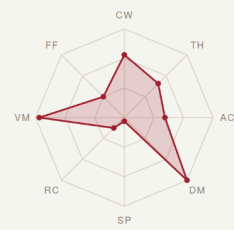
HABITAT ATLAS

Stakeholder Priorities — the same building, weighed differently by each group.

RANK-WEIGHTED
PRIORITY ANALYSIS

Builder / Installer / Manufacturer

n = 7 · 37 stickers



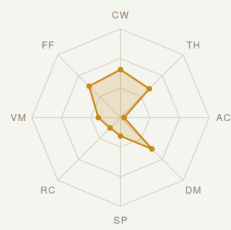
TOP SUBCATEGORIES



Durability and craft lead; values close behind.

Practitioner / Consultant

n = 6 · 26 stickers



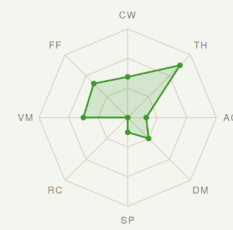
TOP SUBCATEGORIES



The flattest lens — thermal edges a wide field.

Educator / Researcher

n = 5 · 24 stickers



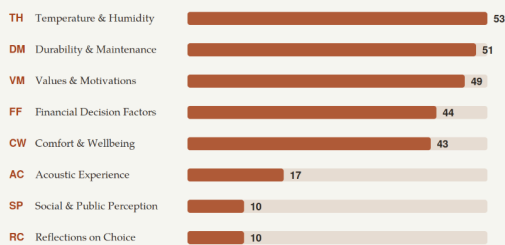
TOP SUBCATEGORIES



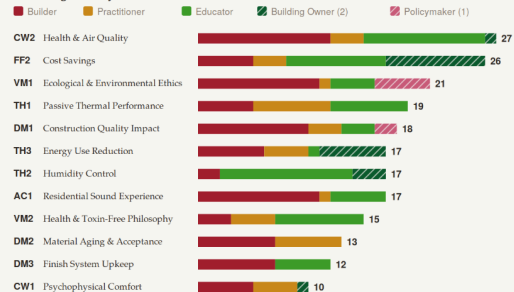
Health and humidity rise to the top.

OVERALL — 101 stickers · 5 stakeholder types · 25 subcategories

Categories by priority



Subcategories by stakeholder



Chosen by no one: AC2 Venue Acoustics · SP3 Public Attention · RC2 Hindsight & Lessons Learned

Method: priority score = sum of ranks (1 low – 5 top) each group placed on a subcategory. Ranks extracted from the wheel photograph; treat scores as ±1–2. Small samples (n=5–7) — read as tiers, not exact order. Hatched segments = a 1-2 person role (individual ballots, not a bloc).

Second Activity - results and findings

Two patterns are worth carrying forward. The comments kept returning to people rather than the material: the worry was rarely whether hempcrete performs, but whether the installer, tradesperson, architect, or inspector knows how to work with it, which echoes the occupant interviews, where finish and installation quality, not the substrate, shaped the lived experience. Several participants, the policymaker among them, asked outright for a benchmark they could cite, which is what the Atlas is designed to be. And the seats that never filled are a result in themselves: with no developer or financier in either activity, the money side of the decision, the part most central to de-risking, went unrepresented.

WHY IT MATTERS FOR THE FIELD

The field has defaulted to cost and energy proxies because occupant evidence in a usable form did not exist. This session produced a first read on what each stakeholder type actually weights, which is the input the next version of the Habitat Atlas needs in order to show each role the evidence that matters most to it first.



BBMC Authored by MASS.

2026 Summit Proceedings

Several results translate directly into the build. The cost finding argues for weighting priorities by intensity rather than counting them, so a minority-but-decisive concern is not lost in the average. The written gaps give a concrete shortlist of subcategories to test in the next round of interviews: lifetime cost to occupy, the disruption of building and moving in, end-of-life, and accessibility. The three dead spokes mark where the framework currently asks about things this audience does not weigh, and where the next interview round can ease off.

Beyond the tool, the session reframes the adoption barrier itself. The recurring worry about whether the trades can execute, set against occupant evidence that the material performs reliably, points the field toward workforce knowledge and citable benchmarks rather than further material testing. The Atlas is built to be the benchmark that participants kept asking for.

The absent seats tell organizers and the industry who needs to be in the room next. Bringing developers and financiers into a future round would close the one gap that most limits these findings. As pre-validation of the Atlas decision layer, the session also strengthens the case for the federal and state funding that would scale this work nationally.

WHAT THIS CAN AND CANNOT TELL US

This was a small, self-selected, single-event sample, and it should be read as direction rather than representation. Group sizes of five to seven people are too small to support a significance test, so the sheets are best read as tiers, with a clear top, a softer middle, and a tail of one-off picks. The wheel filled over the course of the day, so later participants placed their stickers in view of earlier ones, and the picture accumulated as a group conversation rather than a set of independent ballots. Holding everyone to five selections shows what people value without fully showing how much, which is part of why the wheel's intensity reading complements the count. I treat these results as pre-validation for the Atlas decision layer and as a way to sharpen the next round of occupant interviews, not as findings that stand on their own.