



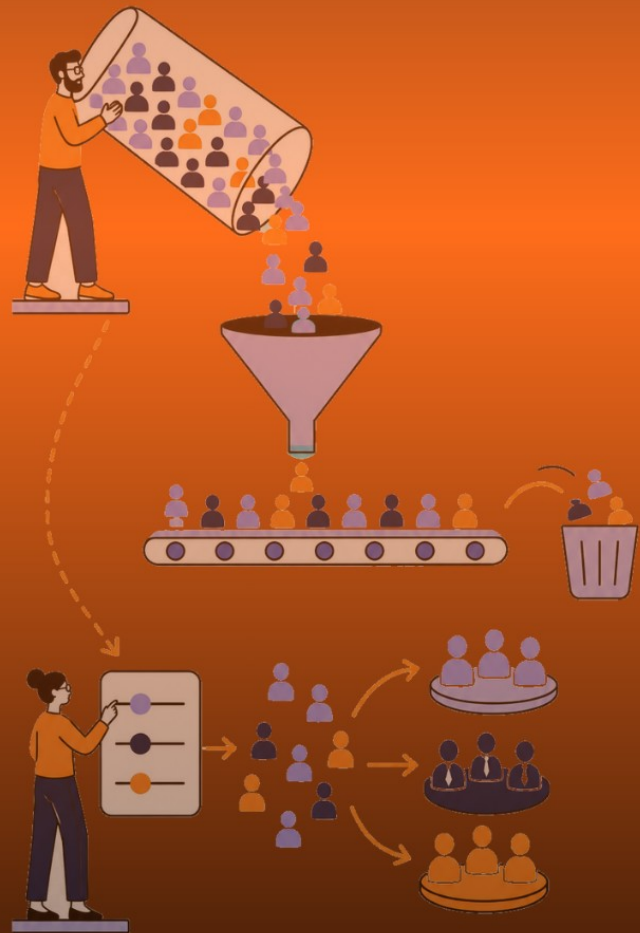
THE

DECISION SEQUENCE

MARKETING BUILT AROUND HOW PEOPLE DECIDE

A list is not a segment

How to
score a list
into
segments



Building the rubric everything else assumes you already have



THEDECISIONSEQUENCE.COM

A J. WORDEN COMPANY · JWORDEN.COM

How to turn a list into segments

Building the rubric everything else assumes you already have

Gates. Relevance: secondary. Credibility: no. Proof: secondary. Safety: no. **Access: primary.**

Score every prospect on three things you can see from outside the account: how much validation they need before engaging, who holds the decision, and what they've adopted early in the past. Psychology describes a segment. Behavior places one.

The symptom

The model went over well. Five segments, clean definitions, most of the room nodding, two people already naming accounts out loud.

Then Monday.

The list opens and it's still sorted by revenue, because revenue is a column that exists. Somebody adds a Segment field to the CRM. Three weeks later it's forty percent blank, and the entries that did get filled in came down to whether an account felt like an early adopter to whoever was typing that afternoon.

Nobody in this story is being lazy. Ask the person filling in the field to describe an Early Mainstream buyer and they'll do it well, in detail, with a fair amount of feeling about the last one they dealt with. Ask which of the two hundred names in front of them is one and the room gets quiet.

Why it keeps happening

Segment definitions get written from the inside out. They describe what each type is like: how they hold risk, what reassures them, what makes them stall a decision for six months. That's the right way to write them if the goal is writing *to* somebody. It does nothing for finding them.

Placing a real prospect requires observable evidence, and psychology isn't observable. You can't look at a company and see caution.

So the model stays a diagram on slide fourteen, the list stays sorted by revenue, and both of those are reasonable responses.

I'll take this one. We defined five segments, described the psychology of each in detail, wrote the sequencing guidance that assumes placement has already happened (snippet 24) and the persona work that assumes you know which segment you're writing into (snippet 25), and never said what to look for. The criteria weren't buried in an appendix somewhere. They didn't exist. Every firm that couldn't apply the model was responding correctly to a model that hadn't finished.

What the buyer is doing at the Access gate

Deciding whether this is worth the next twenty minutes. That decision turns almost entirely on whether the approach matches where they already are, and they have no idea you've scored them.

Approach an Early Mainstream buyer the way you'd approach an Anchor and you're asking somebody to go first who has no intention of going first. The ask reads as exposure. Run the committee reassurance package at an Anchor and you've handed a person who trusts their own read a stack of other people's opinions, which reads as slow.

A wrong placement costs more than no placement at all. Un-placed outreach is obviously generic and gets ignored cheaply, and the contact survives it. Wrong placement fires on schedule, with confidence, in the right week, with the wrong opening line, and it spends something you only get once.

The correction

Build the rubric. Three dimensions, all of them checkable from outside.

How much validation they require before engaging. None, and they'll go first: Anchor. Validation from a respected first mover: Early Adopter. Peer adoption evidence and committee consensus: Early Mainstream. Broad consensus and visible discomfort at being last: Late Mainstream. Nothing resolves it: Laggard.

Who holds the decision. One person who trusts their own read points to Anchor or Innovator. A committee with documentation requirements points to Early Mainstream. A committee plus an incumbent relationship points to Late Mainstream.

What they've done before. Observable history rather than stated preference. Have they hired an unfamiliar firm in the last three years. Have they adopted anything before their peer group did. Have they changed a long-standing provider. Past adoption behavior is the single most reliable signal of the three, and it's publicly checkable more often than firms expect.

Where the evidence lives: their own published material, who they've hired, how they behave at conferences and on panels, whether they speak about things before their peers do, and what your own team already observed in prior conversations and never wrote down anywhere.

Then run it. Top fifty names, one sitting, with the two or three people who have talked to those accounts. Score fast and accept that it's provisional.

The discipline sits in what gets recorded. Put the evidence beside the score. A bare score can't be revised, because there's nothing to revise it against, and six weeks later nobody remembers whether the B was research or a mood.

Name the review cadence. Name the rule for moving somebody, which is new observed behavior. An account doesn't get promoted because the quarter is closing and somebody feels good about the last call.

Before and after

Before. Tier A: revenue over \$50M. Tier B: \$10M to \$50M. Same five-touch sequence, staggered by tier.

After. Account 12, regional distributor. Validation: needs peer evidence, declined to pilot in 2023 until two references were produced. Decision: committee of five, procurement template on file. History: same

logistics provider fourteen years, no unfamiliar hires found. Early Mainstream, leaning Late. Evidence: their own RFP language, plus the call notes from March. Open with the named peer.

When this doesn't apply

Thirty names and two people who have spoken to all thirty. The rubric is overhead. Say the placements out loud, write them down, move on. It earns its keep when the list is longer than anybody's memory of it.

Run this on your own material

Take five names off the current list. Place each one using only what you can observe: what they've published, who they've hired, how the last conversation went, what they did before their peers did it. Give it fifteen minutes total.

Then note which dimension you couldn't answer. Whichever one goes blank first is the one your research process doesn't cover, and it'll go blank on the other forty-five too. That's a fixable problem and a cheap thing to learn on a Tuesday.

Signals you've cleared it

Advance when every name in the top fifty carries a placement and a line of evidence underneath it. When somebody argues with a placement by citing something they saw. When the opening line of the outreach changes depending on the score, which is the only reason the score exists.

Hold when the segment field is populated and the evidence column is empty. Hold when the placements cluster: forty of fifty scored Early Mainstream usually means the scorer defaulted to the middle rather than looked. Hold when nobody can say what would move an account from one segment to the next, because a rubric nobody can revise turns into a label, and labels age badly.

Where this connects

Story from the Field · *The Decision You Could Survive*, on the buyer whose real question is defensibility.

Research Library · Entry 05, status quo bias. Entry 08, authority and social proof.

Next move · Go to *How to sequence outreach*.