



Attendee Segmentation: Research & System Proposal

How Flowz should model lifecycle, tiers, badges and scoring - grounded in a nine-track research sweep across HubSpot, Salesforce, Klaviyo, Zoho, ActiveCampaign, Braze, Iterable, CleverTap, the event-ticketing industry, canonical RFM practice, loyalty-program naming psychology, and the documented failure modes of scoring systems.

TO REVIEW · INTERNAL

August 2026 · Companion deck: flowz-hub.pages.dev/tiers · Prepared for Elad &

Aviv

• Contents

Part 1	Executive summary
Part 2	The proposed system: three fields, one engine
Part 3	Tier naming systems (five candidates)
Part 4	Architecture options considered
Part 5	Client editability, defaults and guardrails
Deep dive 1	HubSpot
Deep dive 2	Salesforce
Deep dive 3	Klaviyo
Deep dive 4	Zoho CRM + ActiveCampaign
Deep dive 5	Braze · Iterable · CleverTap
Deep dive 6	Event-Industry Platforms
Deep dive 7	The RFM(E) Model
Deep dive 8	Loyalty-Tier Naming
Deep dive 9	Failure Modes

1 Executive summary

Flowz today models attendees on three unconnected verticals: lifecycle stage, membership tier and freeform tags. The research verdict is that the three-field decomposition itself is industry-correct - every major platform keeps a temporal state, an earned status and descriptive labels apart. What is missing is the connective tissue: one data engine underneath, hard boundaries between the axes, and names a stranger can rank instantly.

The proposal in one paragraph. Every attendee action lands in an append-only ledger. Three engines read it: Recency + Engagement drive **Lifecycle** (temporal state, counted in missed editions rather than calendar days), Frequency + Monetary drive **Tier** (earned status: events attended AND lifetime value over a rolling window), and pattern rules drive **Badges** (Early Bird, Group Buyer, Big Spender, Veteran). The boundary rule: rankable signals feed tiers, temporal states are lifecycle, behavior patterns are badges. Nothing can contradict, and any threshold a client edits later recomputes cleanly from history.

The strongest findings

- **Score $\neq$ tier $\neq$ stage, everywhere.** HubSpot, Salesforce, Zoho and ActiveCampaign all keep them separate; platforms that fuse them (Klaviyo, CleverTap) lock the vocabulary and hide the math.
- **Scores are plumbing, names are the product.** No event platform exposes a raw 0-100 fan score. Everything surfaces as 4-7 named states.
- **Calendar recency is Flowz's biggest exposure.** "No purchase in 365 days" marks every annual-festival fan lapsed in month 11. Recency must be event-relative: missed editions of a series the attendee used to attend.
- **Tier physics matter.** Upgrade instantly, downgrade only at season end with notice and a grace window - demoted customers measurably end up less loyal than never-promoted ones.
- **Clients edit definitions, never the math.** Sentence-style rules, histogram sliders with live counts, impact previews before saving, defaults seeded from the client's own history.

2 The proposed system: three fields, one engine

The ledger

An append-only record of value events per attendee: ticket purchases (with edition number), room / F&B / merch spend, message opens, clicks and replies, WhatsApp and support touches, group size, buy timing, referrals. Balances are aggregates over the ledger (the Salesforce LoyaltyLedger pattern), which is what makes later threshold edits safe: tiers recompute from history instead of being patched in place.

The three fields

FIELD	DEFINITION	DRIVEN BY
Lifecycle	Temporal state: New → Returning → Regular → Cooling → Lapsed. Recency counted in missed editions, never calendar days. Engagement splits Cooling into "reachable" (silent wallet, still opens everything) vs "silent".	Recency + Engagement
Tier	Earned status. Qualification = events attended AND LTV over a rolling 24-month window, so a single VIP-table purchase cannot buy the top tier. Upgrades apply instantly; downgrades only at season end with notice and a grace window; tier freezes in the producer's off-season.	Frequency + Monetary
Badges	Behavior patterns, auto-computed from editable rules, in categories: Buying (Early Bird, Last-Minute, Group Buyer), Spend (Big Spender, Rooms, Merch), Engage (Responsive, Advocate), Tenure (Veteran, Founding Fan). A separate freeform manual layer (Artist, Helper) with a Tag Manager showing usage counts.	Pattern rules

The five logic questions

SIGNAL	LIVES IN	WHY
Events attended	Tier input	Rankable. Rolling 24 months, 3-4 bins - quintiles over "attended 1, 2 or 3 events" is fake precision.
Rooms / F&B revenue	Tier input	The Monetary sub-score. Events AND LTV together mean a rooms-always buyer with 2 events can outrank a GA-only buyer with 4.
Group buyers	Badge	Not rankable. Their money already counts in Monetary; the badge captures the pattern for group offers and referral flows.
Early / late buyers	Badge	Gold for pricing flows, meaningless as rank. Early Bird = first-window buyer 2+ times.
Engagement score	Lifecycle modifier	Ratio with decay (share of messages engaged, events attended vs opportunity). Splits identical recency states; never feeds the tier, because blended scores hide the actionable driver.

3 Tier naming systems: five candidates

Naming rules from the research: a stranger must rank the ladder in under 2 seconds; 1-2 words per name; the top tier may break the naming pattern to feel "beyond the ladder" (Rouge, Centurion, Reserve); no word may appear on two axes; keep the top tier small and the bottom tier subordinate. Architecture: stable internal keys (tier_1..tier_4) with client-editable **name skins** - the Delta model of fixed names over repriced numbers.

SYSTEM	LADDER	CHARACTER
WAVES Flowz brand	Ripple → Wave → Surge → Tsunami	Energy magnitude in water terms. The crowd version - bass, parties, festivals.
WATERS Flowz brand	Drop → Stream → River → Ocean	Scale of water. The universal version - needs zero explanation in any vertical.
BACKSTAGE Cool · safest	GA → Front Row → Backstage → All-Access	Venue proximity physically encodes rank; industry vernacular every promoter decodes instantly.
FIRE Cooler · genre skin	Spark → Glow → Blaze → Inferno	Heat self-ranks. Born for EDM and bass brands; deliberately wrong for yoga - that is what skins are for.
APEX Coolest · dynamic	Rookie → Riser → Elite → Apex	The top tier is not a threshold: it is the top 5% of this producer's audience, live. Equal exclusivity at any audience size, zero tuning. No CRM ships this as a first-class construct.

Recommendation. Backstage as the default skin, Waters for wellness clients, and the Apex dynamic top-% as a differentiator feature regardless of which names ship. Lifecycle rename proposal: New → Returning → Regular → Cooling → Lapsed ("Rising" and "Core" need a legend; "Returning" does not).

4 Architecture options considered

APPROACH	ASSESSMENT	VERDICT
A · Three fields + one engine	Lifecycle, Tier and Badges as visible fields; recency, frequency, monetary and engagement engines underneath, all reading one ledger.	RECOMMENDED
B · Fused RFM grid	One set of named cells (Champions, At Risk, Can't Lose Them...) as in Klaviyo and CleverTap. Loses "Gold member going cold", the highest-value insight. The LTV dashboard's lifecycle × tier cross-view already renders this grid, so the consolidation happens in the view layer.	KEEP AS VIEW
C · Fan Score 0-100	One blended number per attendee. Composites hide drivers; no event platform exposes one; roughly 40% of sales teams trust opaque scores.	REJECTED
D · Fan-facing points program	Points, redemptions, leaderboards (Tixr Rewards style). A different product for a later phase; its ledger and currency-split lessons apply now.	LATER PHASE
E · Two-axis grade	HubSpot-style A1-C3: value letter × recency number. Power-user shorthand for grid cells, not a primary construct.	CELL LABELS

5 Client editability, defaults and guardrails

The config surface

- **Sentence rules:** "Backstage = attended [4]+ events AND spent [≥2,500]+ over the last [24] months" - every bracket editable inline (the ActiveCampaign pattern).
- **Histogram + drag handles:** the producer sees the live distribution of their own attendees and drags band boundaries; counts update in real time; one click materializes a band as a segment (Braze).
- **Impact preview before save:** "This change moves 1,240 attendees Backstage → Front Row" (Klaviyo).
- **Seeded defaults:** percentiles over the client's own history propose the starting thresholds (the HubSpot AI-draft pattern). Nobody starts from a blank rule builder.
- **Pro plan:** a Flowz representative tunes thresholds with the client - the consultant-assisted model Zoho's partner ecosystem proves demand for, built into the plan instead.

Shipping defaults

SETTING	DEFAULT
Frequent / core	3+ events per rolling 24 months
Top tier size	≤ 5-10% of the audience
Lapsed	2 missed editions of a series previously attended
New	first purchase, first 90 days
Tier window	rolling 24 months; season-end downgrade with grace

Guardrails from the failure research

- Hard caps: at most 5 lifecycle stages and 4 tiers. One segment = one default action; if two segments would trigger the same campaign, merge them.
- No weight sliders, ever - weights are meeting-room guesses nobody backtests.
- Dependency guard: "used in 3 flows" blocks deletion (the HubSpot "Used In" pattern).
- Explainability: every state shows its fired rule ("Cooling ↓ - 1 missed edition, opens down 60%").
- Vocabulary rule: no word appears on two axes.
- Validation nudge: periodically compare definitions against outcomes ("62% of Cooling actually lapsed").

D1 HubSpot

HubSpot (Marketing Hub / Sales Hub, incl. the 2025 lead-scoring engine replacing legacy score properties)

ARCHITECTURE

HubSpot keeps FOUR separate, orthogonal constructs and never merges them. (1) Lifecycle Stage: a single-value, pipeline-like property shared by contacts AND companies – 8 ordered defaults (Subscriber, Lead, Marketing Qualified Lead, Sales Qualified Lead, Opportunity, Customer, Evangelist, Other). It answers "where in the journey." Built-in automation is forward-only: HubSpot can auto-set Opportunity when a deal is associated and Customer on closed-won, and default automation never moves a record backward. A secondary property, Lead Status (New, Open, In Progress, Open Deal, Unqualified, Attempted to Contact, Connected, Bad Timing), acts as sub-states within the sales-qualified phase. Each stage auto-generates calculated properties ("latest time in stage", "cumulative time in stage" – Pro/Enterprise) for funnel reporting. (2) Scoring: numeric, deliberately separate from lifecycle. The new engine splits Fit score (static demographics/firmographics: job title, company size, revenue, region) from Engagement score (behavior: page views, form submissions, email clicks, meetings), with an optional Combined score that writes three properties (total, engagement-only, fit-only). Scores DRIVE lifecycle transitions via workflows (score $\geq$ threshold $\rightarrow$ set stage to MQL), not the other way around. (3) Lists/Segments: Active lists (records join/leave automatically as criteria change) vs Static lists (snapshot at save time). Up to 250 filters per list incl. 60 associated-object filters, AND/OR groups. Lists consume the other constructs (filter on lifecycle stage, score, properties) and feed emails, workflows, ads, reports. (4) Tags: there are NO native contact tags. The tag concept is covered by custom properties (multi-checkbox/dropdown) plus active lists; the only literal "tags" are color-coded object tags on deal/ticket/lead BOARD views – rule-based labels (e.g. green "Large deal" tag = Amount > \$10,000), usable afterwards in views, reports, workflows, segments. Net: stage = journey position (one value, ordered), score = continuous measure with banded labels, lists = derived populations, tags = rule-based visual flags.

VERBATIM NAMING

- Lifecycle stages (verbatim, in order): Subscriber, Lead, Marketing Qualified Lead, Sales Qualified Lead, Opportunity, Customer, Evangelist, Other
- Lead Status values: New, Open, In Progress, Open Deal, Unqualified, Attempted to Contact, Connected, Bad Timing
- Score types: Engagement score, Fit score, Combined score (legacy property was literally named 'HubSpot score' – frozen Aug 31, 2025)
- Score threshold labels: High, Medium, Low (single scores); combined scores get a 9-cell grade grid A1, A2, A3, B1, B2, B3, C1, C2, C3 – letter = fit quality, number = engagement level, 'A1 is your ideal lead'
- Segments: Active list, Static list
- Board labels: Deal Tags, Ticket Tags, Lead Tags (color-coded object tags)

WHAT CUSTOMERS CAN EDIT

Lifecycle stages: editable at Settings > Objects > Contacts > Lifecycle Stage tab, Super Admin only, one shared stage set for contacts+companies. You can add custom stages (e.g. 'Onboarded', 'Lost Customer'),

rename inline, reorder via drag handle, and delete – BUT deletion is blocked while any record holds the value or the stage is referenced in lists, workflows, reports, saved views, ads, chatflows, or as a connected-app default; a 'Used In' column shows a clickable count of references. Third-party sources conflict on exact lock rules (several state Subscriber and Customer cannot be renamed; the official KB just says defaults can be 'customized'), so treat a small set of anchor stages as effectively locked. Scoring: everything is editable – criteria, add/subtract points per rule (decimals allowed, negatives allowed: Student/HR -5, Do-Not-Target list -30), per-group point caps, overall score cap (100 points Professional / 500 Enterprise), frequency operators (Exactly, Between, At least, At most) and caps ('CTA click +10, capped at 3'), time windows (7/14/30 days) OR linear score decay at 1/3/6/12-month intervals (cannot combine both in one group), aggregation for associated-company criteria (Sum/Average/Minimum/Maximum). Threshold bands for High/Medium/Low and the A–C / 1–3 grid are user-set ranges (example: A = 38–50 of a 50-pt fit group, 1 = 35–50 of a 50-pt engagement group). Limits: 5 scores per object on Professional, 50 on Enterprise. Engagement scores can be reset via workflow (closed-lost, 180-day inactivity, unsubscribe); 'You can't reset Fit.' Lists: fully custom criteria, all plans. Fixed: the lifecycle property is single-value; default automation only moves stages forward; combined deal scores are Sales Hub only; AI scoring is Marketing Hub Enterprise only.

CONFIG UX

Two very different editing surfaces, matched to complexity. Lifecycle stage config is a dead-simple ordered-list editor (like a pipeline editor): rename in place, drag-handle to reorder, hover-Delete with dependency guard ('Used In' count links to every asset referencing the stage), plus toggles for automation (default stage for new records, auto-update on deal association/close, defaults for records synced from connected apps). No rule builder – stages are labels with semantics, and transition logic lives in workflows. Scoring config is a dedicated builder (Marketing > Lead Scoring): name the score → choose object and type (engagement / fit / combined) → add criteria organized into groups → per-rule points with add/subtract, frequency operators, and time windows → per-group max points → overall max → optional decay per group → 'Review and turn on'. Threshold sliders define the High/Medium/Low bands and the A1–C3 grid cells. Enterprise adds AI-drafted scores: you pick a lifecycle transition to optimize for (e.g. 'Subscriber' to 'Sales Qualified Lead') and a timeframe (e.g. 90 days), HubSpot trains on your historical contacts (up to an hour), proposes criteria/weights, and you edit before activating. Records show a score-history card explaining point changes. Lists use a filter-builder with AND/OR groups. Everything sensitive is admin-gated (Super Admin for lifecycle edits) and consultant-friendly, but designed so a non-technical marketer can read a record and instantly parse 'MQL, score 72, grade A2'.

TAKEAWAYS FOR FLOWZ

- Keep lifecycle and tier/score as two orthogonal dimensions and let score DRIVE stage transitions via rules – HubSpot never collapses them; lifecycle is a single-value ordered pipeline, score is a capped number. Flowz's lifecycle × tier model is structurally right; the redesign should formalize 'tier = f(score)' and 'stage transitions triggered by score/recency rules' instead of three parallel unrelated systems.
- Raw points are unreadable; HubSpot always translates them into small labeled bands: High/Medium/Low, or the two-axis A1–C3 grid (fit letter × engagement number). For Flowz, a two-axis grade like spend-letter × attendance-number (e.g. 'A2 = big spender, moderately frequent') is instantly explainable to a festival producer, and the band THRESHOLDS are what clients edit – not the formula, which is v2 territory.
- Ship opinionated defaults with guarded editing: rename/reorder/add stages in a simple list editor, admin-only, with a 'Used In' dependency count that blocks deletion while records or automations reference a stage, and forward-only automatic transitions. This is exactly the safety model Flowz needs so B2B clients can customize tiers/stages later without breaking flows.
- Copy the capping-and-decay toolkit: overall score cap (HubSpot: 100/500 pts), per-group caps, per-rule frequency caps ('+10 per CTA click, capped at 3'), and linear decay at 1/3/6/12 months. For attendees, decay IS the At-risk/Lapsed mechanic – derive 'At-risk' from engagement-score decay rather than maintaining separate hand-written recency rules.
- Don't build tags as a freeform first-class system. HubSpot has no contact tags at all – 'tags' are either structured properties, dynamic list membership, or rule-based colored board labels ('Large deal' = Amount > \$10,000). Flowz tags like Spender/GroupBuyer should be rule-based derived flags (editable threshold, e.g. Spender = LTV > X) that are reusable in segments and automations, not manually applied stickers.
- Offer an AI-drafted starting model as the premium path: HubSpot Enterprise trains scoring on 'who reached stage Y within 90 days' and proposes weights the admin edits before turning on. For Flowz, seeding a client's tier thresholds from their own historical attendance/LTV distribution (then letting them tweak) beats asking producers to invent point values from scratch.

D2 Salesforce

Salesforce (Loyalty Management + Sales Cloud Einstein Lead Scoring + Account Engagement/Pardot Scoring & Grading + Marketing Cloud Einstein Engagement Scoring)

ARCHITECTURE

Salesforce deliberately keeps THREE orthogonal constructs and never merges them. (1) VALUE TIERS live in Loyalty Management (separate Industries license): LoyaltyProgram -> LoyaltyTierGroup -> LoyaltyTier objects. A tier group defines the time mechanics: Tier Model = "Fixed" (qualifying period starts/ends on preset dates, e.g. points reset every Dec 31) or "Anniversary" (period runs from each member's enrollment date to their anniversary), plus Tier Period (how long earned status is guaranteed, typically 1 year) and Extend Expiration (defer downgrade to month-end / reset date / anniversary). Each tier is just: Name, SequenceNumber (rank), MinimumEligibleBalance / MaximumEligibleBalance (threshold band of qualifying points), Color. Critically, every program has TWO currencies: a qualifying currency (e.g. "Tier Points") that only drives tier status and resets to zero each qualifying period, and a non-qualifying currency (e.g. "Regular Points"/"Member Points") that members spend on rewards and expires separately (fixed 3-year or activity-based). Spending rewards can never demote you. Upgrades are assessed in real time when a qualifying activity posts; downgrades happen only at tier-period end via batch Data Processing Engine jobs (verbatim job: "Reset Qualifying Points"). Every point movement is an append-only LoyaltyLedger row aggregated into balances – no overwriting a points field. Benefits (free delivery, early access, birthday gift, welcome-to-tier) attach per tier with a priority rank. (2) SCORES are a separate layer with two philosophies: Einstein Lead Score is ML, 0-100 (opportunities 1-99), trained on your past conversions (needs 1,000 leads/180 days, 120 conversions, else falls back to an anonymized global model), shown with top positive/negative factors for explainability; Pardot/Account Engagement Score is rules-based (activity -> points table, fully editable, retroactive), paired with a Grade (A-F in thirds: C, C+, B-, B...) measuring profile FIT vs ideal customer, plus an optional Einstein Behavior Score (0-100, recency/frequency-weighted ML). The canonical pattern is score (behavior) x grade (fit) as a 2-axis matrix. (3) LIFECYCLE/BEHAVIOR SEGMENTS live in Marketing Cloud Einstein Engagement Scoring: per-subscriber 7-day likelihood predictions (open, click, web conversion, retention) written to the MC_Einstein_Predictive_Scores data extension, collapsed into exactly four personas from an open-likelihood x click-likelihood quadrant: Loyalists (high/high), Window Shoppers (open but don't click), Selective Subscribers (rarely open but click when they do), Winback/Dormant (low/low). Personas feed Journey Builder "Scoring Split" branches. So: tier = earned, sticky, customer-facing status; persona = computed, volatile, internal behavior state; score = continuous prioritization number. Flowz's lifecycle (At-risk/Lapsed) maps to the persona layer, Bronze->Platinum to the tier layer, tags to score/grade inputs.

VERBATIM NAMING

- Loyalty tiers (Trailhead 'LevelUp' program): Silver (0 pts), Gold (1,000 pts), Platinum (5,000 pts)
- Loyalty tiers (Trailhead 'Cloud Kicks Inner Circle' program): Silver, Gold, Diamond, Platinum
- Program containers are branded, tiers are generic: 'Cloud Kicks Inner Circle', 'LevelUp'
- Currencies: 'Tier Points' (qualifying), 'Regular Points' / 'Member Points' (non-qualifying)
- Tier models: 'Fixed' and 'Anniversary'
- Engagement personas: 'Loyalists', 'Window Shoppers', 'Selective Subscribers', 'Winback/Dormant'
- Pardot Grades: A to F in increments of thirds (C, C+, B-, B, B+...)
- Scores: 'Einstein Lead Score' (0-100), 'Einstein Behavior Score' (0-100), 'Pardot Score' (unbounded rules-based)
- LoyaltyTier object fields: Name, SequenceNumber, MinimumEligibleBalance, MaximumEligibleBalance, Color, EligibleEnrollmentType (Free/Paid/SegmentBased/TierBased/Trial)
- Batch processes: 'Reset Qualifying Points', 'Expire Activity-Based Non-Qualifying Points'

WHAT CUSTOMERS CAN EDIT

Tiers: fully customizable – tier count (2+, typically 3-4), names, colors, threshold values (MinimumEligibleBalance per tier), tier model (Fixed vs Anniversary), qualifying-period reset date/frequency, tier-period length, downgrade-deferral rules, and which benefits attach to which tier. Rules-based scoring (Pardot): fully editable – admins change per-activity point values in a settings table (Pardot Settings -> Automation Settings -> Scoring); edits apply retroactively account-wide; score decay is NOT built in (must be automated manually). Grading: rules-based, marketer-defined criteria. ML scores (Einstein Lead Score, Behavior Score, Engagement Scoring): essentially NOT customizable – admins can only (a) split leads into up to 35 segments with up to 100 field filters each, giving each segment its own model, and (b) exclude fields from the model. No weight sliders anywhere. Engagement personas: the four names, definitions, and thresholds are fixed by Salesforce – customers cannot rename or re-threshold them.

CONFIG UX

Loyalty setup is a guided admin wizard, not sliders: create program -> create tier group (pick tier model, qualifying period, reset frequency) -> define two currencies -> add tiers as an ordered list of name + minimum-points rows -> attach benefits per tier with priority -> build accrual/redemption rules ('Rule-Based Tier Assessment' is a toggle; buy shoes = instant points, post review = next-day points; redeem 100 points = \$10 off). It's an admin/consultant product (own license, Industries family); end business users see dashboards, not config. Pardot scoring UX is the pattern most relevant to Flowz: a plain table of activities with editable numeric point values, one screen, with an explicit warning that changes recompute history retroactively. Einstein ML surfaces are read-only dashboards; trust is built by always showing top positive and negative factors next to each score rather than letting users edit the model. Marketing Cloud personas surface as a dashboard plus data-extension fields usable in drag-and-drop Journey Builder splits – segmentation as a routing primitive, not a report.

TAKEAWAYS FOR FLOWZ

- Keep lifecycle and tier as two separate dimensions and give them opposite physics: tiers upgrade instantly when a threshold is crossed but only downgrade at period end (e.g. 12 months of guaranteed status); lifecycle stages (At-risk/Lapsed) recompute freely. This asymmetry is what makes tiers feel like earned status instead of a volatile label – Salesforce hard-codes it via Tier Period + batch downgrade.
- Split the tier-driving metric from any spendable/reward metric now, even if rewards come later: Salesforce's qualifying vs non-qualifying currency split exists so redemption never demotes a member and so tier thresholds stay a single clean number. For Flowz: define one 'Tier Points' style qualifying metric (events attended + LTV formula) and keep it isolated.
- A tier is just four fields – Name, rank order, minimum threshold, color – plus group-level time settings (reset model: fixed calendar date vs member anniversary; period length). That's the entire customer-editable surface Flowz needs: an ordered list of rows with editable names and one threshold number each, no weight sliders.
- Naming: Salesforce's own B2C demos never invent exotic tier names – they use the metal ladder (Silver/Gold/Platinum, sometimes Bronze or Diamond) because clients parse it instantly, and put the brand personality in the program container name ('Inner Circle'). For lifecycle stages, copy the persona pattern: plain-English behavioral labels like Loyalists / Window Shoppers / Winback-Dormant beat abstract terms like 'Core'.
- If Flowz adds scoring, choose the Pardot pattern, not the Einstein pattern: a transparent activity->points table that B2B clients can edit (with retroactive recompute), optionally as a 2-axis model (behavior score x fit grade A-F). If any ML score is added, make it read-only and always show top positive/negative contributing factors; Salesforce lets admins customize ML only by segmenting models and excluding fields, never by editing weights.
- Store attendee value events as an append-only ledger with aggregated balances (Salesforce's LoyaltyLedger pattern): it is what makes 'client edits a threshold later' safe, because tiers can be recomputed from history instead of patched in place.

D3 Klaviyo

Klaviyo

ARCHITECTURE

Klaviyo has NO first-class "lifecycle stage" or "tier" object. It has four constructs that layer: (1) **PROFILES** carry properties – both custom properties and system-computed ones. (2) **PREDICTIVE ANALYTICS** are read-only computed profile properties (Historic CLV, Predicted CLV = predicted spend next 365 days, Total CLV = historic+predicted, Churn Risk Prediction = 0-1 probability of not purchasing again, Average Time Between Orders, Expected Date of Next Order, Predicted Gender). Models retrain at least weekly; they only unlock once the account has 500+ customers who placed an order, an ecommerce integration or API order events, 180+ days of order history with an order in the last 30 days, and some customers with 3+ orders. (3) **RFM ANALYSIS** is a report in Intelligence/Marketing Analytics that scores every customer 1-3 on Recency, Frequency, Monetary (percentile-based by default: R3 = purchase within 180 days, R2 = within 365, R1 = older; F and M = top/middle/bottom thirds), then maps the 27 score combinations onto exactly 6 named groups (e.g. Champions = 333/332/323; Inactive = 111/112/113/121). The group – not the raw score – is written nightly onto each profile as three properties: "Current RFM group", "Previous RFM group", "RFM group last changed". (4) **SEGMENTS** are the single consumption layer: dynamic, condition-based, near-real-time groups built with plain-English rows like "Properties about someone > Current RFM group > equals > Champions" or "Predictive analytics about someone > Churn Risk Prediction > over 0.7", combinable with And/Or. Lists are static opt-in collections (grow/shrink only by subscribe/unsubscribe); Tags are account-level organizational labels for lists/segments/flows/campaigns and explicitly CANNOT group profiles – profile grouping is done via custom properties + segments. So: scores are invisible plumbing, the named group is the customer-facing unit, and everything funnels into one segment builder. Note the contrast with Flowz: Klaviyo collapses lifecycle AND value-tier into the single RFM group label (Champions is simultaneously "high value" and "active"), whereas Flowz currently keeps lifecycle and tier as two orthogonal dimensions.

VERBATIM NAMING

- RFM groups (verbatim, fixed set of 6): Champions, Loyal, Recent, Needs attention, At risk, Inactive
- RFM profile properties: Current RFM group, Previous RFM group, RFM group last changed
- Predictive metrics: Historic CLV, Predicted CLV, Total CLV, Churn Risk Prediction, Average Time Between Orders, Expected Date of Next Order, Predicted Gender
- Segment condition categories: 'Properties about someone', 'Predictive analytics about someone'
- Audience constructs: Lists (static), Segments (dynamic/condition-based), Tags (object organization only, never profile grouping)
- RFM scores: 1-3 per dimension (not 1-5), e.g. Champions = score combos 333/332/323

WHAT CUSTOMERS CAN EDIT

EDITABLE: (a) RFM score thresholds per dimension via Advanced settings > RFM Details – two modes per dimension: percentile-based (enter any percent 1-100, e.g. "top 75th percentile gets score 3") OR value-based absolute numbers (e.g. "minimum all-time spend \$250 = top Monetary score"; whole numbers only).

A Preview button validates the redistribution before Save; profile properties refresh nightly, so changes hit profiles within 24h. (b) The conversion metric that defines an "order" – defaults to Placed order, switchable to any value-based metric including custom metrics (custom metrics take up to 48h to reflect). (c) Report date range. FIXED: the six group names; the 1-3 score scale; the score-combination-to-group mapping; which properties exist. Predictive analytics has ZERO user tuning – no weights, no model knobs, no threshold editing; it simply appears once the data-volume gates are met and retrains itself weekly. Access to RFM config is role-gated (Owner/Admin/Manager/Analyst only).

CONFIG UX

Report-first, config-buried: RFM lives as a report under Advanced KDP/Intelligence > Customer insights > RFM analysis; the customer sees a dashboard of the 6 groups first, and threshold editing hides behind "Advanced settings > RFM Details" – select Recency/Frequency/Monetary, toggle "Score based on percentiles" or type value thresholds, Preview, Save. It is a settings panel with 3 dimensions and 2 input modes, not a rules engine and not sliders. Predictive analytics has no config surface at all: values just appear on the profile's "Metrics and insights" tab when eligibility is met. All day-to-day targeting happens in the segment builder using readable condition rows (dimension + operator + value, And/Or connectors) – e.g. "Predicted CLV over 200 AND Churn Risk Prediction over 0.7". Transitions are first-class: because Previous RFM group and RFM group last changed exist, customers build flows triggered by "entered At risk" without any extra machinery.

TAKEAWAYS FOR FLOWZ

- Fix the vocabulary, free the math: Klaviyo's 6 group names (Champions/Loyal/Recent/Needs attention/At risk/Inactive) are immutable across all accounts – instant comprehension – while thresholds underneath are editable per account. Flowz should lock a small named stage/tier vocabulary and let producers edit only the numbers.
- Scores are plumbing, names are the product: Klaviyo computes 1-3 scores per dimension (27 combos) but customers only ever see 6 named groups; the raw score is never the UI. If Flowz adds a scoring system, surface it as named segments, not numbers.
- Threshold editing UX = one settings panel with two modes: per-dimension choice of percentile (1-100) or absolute value (e.g. \$250 all-time spend), plus a Preview-before-Save step showing redistribution. That is the entire config surface – copy this instead of building a rule engine for v1.
- Separate deterministic from predictive, and gate predictive by data volume: RFM groups are transparent arithmetic on real behavior; CLV/churn predictions are read-only, un-tunable, and only unlock at 500+ purchasers / 180 days history / 3+ order customers. Flowz can ship editable RFM-style tiers now and add opaque predictions later behind similar data gates.
- One consumption layer: RFM groups and predictions both land as profile properties consumed by a single dynamic segment builder with plain-English conditions; Klaviyo tags deliberately cannot group people – Flowz's attendee tags (EarlyBuyer, Spender...) map to Klaviyo's custom profile properties, and Flowz should keep one unified 'attendee filter' engine rather than three parallel systems.
- Store transitions, not just state: 'Previous RFM group' + 'RFM group last changed' (nightly refresh) are what make 'just dropped into At risk' automations possible. Flowz should persist stage/tier transition history as queryable fields from day one.

D4 Zoho CRM + ActiveCampaign

Zoho CRM (Scoring Rules / Multiple Scoring Rules) + ActiveCampaign (Contact & Deal Scoring, Tags) – combined report, clearly labeled per platform in each field

ARCHITECTURE

ZOHO CRM: Numeric scores are a layer ON TOP of records, not a lifecycle. A "Scoring Rule" is a named per-module object (Leads, Contacts, Accounts, Deals, custom modules; scoped to one layout or "All Layouts"). Each rule has two rule sources: (a) field-based criteria rows (record attributes) and (b) "touchpoint" criteria per integrated channel (Email opened/clicked/bounced/sentiment/intent/keyword, Calls, Campaigns, Webinar registration, Survey responded/visited, Facebook, X/Twitter, Business Messaging, Desk tickets, Backstage purchase/check-in, SalesIQ chat, VOC churn/competitor mentions – options shown depend on what's integrated). Every record then displays SIX derived score categories: Positive Score, Negative Score, Total Score (from record attributes/actions) and Positive/Negative/Total Touchpoint Score (from interactions; touchpoint scores only on Leads and Contacts). Scores can optionally be mapped to real record fields ("Would you like to add score fields to the records?" – consumes numeric custom-field limits, up to 6 score fields per rule) so they become filterable/reportable/workflow-triggerable like any field. There are NO built-in stages or tiers tied to scores – thresholds are the admin's job via views/workflows on the score field. Parallel AI option: "Zia Scores" (health, engagement, follow-up, field attribute, conversion scores) auto-computed instead of manual rules. ACTIVECAMPAIGN: Two score families – "Contact Score" (person quality/engagement) and "Deal Score" (opportunity quality) – plus flat Tags as a third, orthogonal construct. A score is a named container of rules; each rule = segment-builder conditions + points (+/-) + per-rule expiration. Unlimited scores per account, no max score value. Two computation modes with verbatim names: "Static Scoring" (rules on the Contacts > Scoring page, applied once per contact, retroactive) and "Dynamic Scoring" (an automation with trigger "Runs: Multiple Times" + an "Adjust score" action, fires on every occurrence – this is how per-event point accrual works). Tags are explicitly framed as "temporary data": applied manually, by automations, forms, link clicks, integrations, API; they act as automation triggers ("Tag is added"), segment conditions, and conditional-content conditions. Lifecycle-ish states are NOT native objects – ActiveCampaign ships them as a pre-built two-part "Engagement Tagging" automation recipe that adds/removes four tags (Engaged / Recent activity / Disengaged / Inactive) based on time since last email open, click, or site visit, with the time intervals pre-filled and editable.

VERBATIM NAMING

- ZOHO: "Scoring Rules" (Setup > Automation > Scoring Rules)
- ZOHO: "Manual Score" vs "Zia Score" (rule Type selector)
- ZOHO: six record score categories – "Positive Score", "Negative Score", "Total Score", "Positive Touchpoint Score", "Negative Touchpoint Score", "Total Touchpoint Score"
- ZOHO: Zia score types – "health scores", "engagement scores", "follow-up scores", "field attribute scores", "conversion scores"
- ACTIVECAMPAIGN: "Contact Score" and "Deal Score" (chosen in the Add New Score modal)
- ACTIVECAMPAIGN: "Static Scoring" vs "Dynamic Scoring"
- ACTIVECAMPAIGN: "Adjust score" automation action (add / reduce / set specific value / reset to 0) with "Expires" setting
- ACTIVECAMPAIGN: "Score Changes" automation trigger (start automation when score crosses a value, e.g. 75+)
- ACTIVECAMPAIGN: "Tag is added" trigger; "Tag Manager" (merge, bulk edit, per-tag description, shows contact counts + which automations use each tag)
- ACTIVECAMPAIGN: Engagement Tagging recipe tags – "Engaged", "Recent activity", "Disengaged", "Inactive"
- ACTIVECAMPAIGN documented tag naming conventions – "Customer - Camera", "[ACTION] Downloaded whitepaper", "Interest - Product - Cameras", "Visited pricing page"

WHAT CUSTOMERS CAN EDIT

ZOHO – everything is admin-defined, almost nothing is default: rule name, module, layout scope, each criteria row (field + operator + value + Add/Subtract action + free numeric points), and per-channel touchpoint points (e.g. consultant-typical values: +5 email open, +10 click, +5 survey visit, +10 survey response, -50 to -100 for campaign opt-out; example from Zoho docs: "Add 10 points if calls are attended in the last 20 days"). Fixed/guardrailed: the six score categories themselves (can't rename), max 10 scoring rules per layout with max 5 active simultaneously, retro-scoring only touches records created/modified in the last 6 months, needs "Manage Automation" profile permission, channel availability gated by edition (Standard edition: criteria/scoring for Calls channel only; VOC needs Enterprise/Ulimate/CRM Plus/Zoho One 15+ licenses). Rules can be cloned, deactivated (Status toggle to "Inactive" – retains scores, halts recalculation), or deleted (deletes all scores + mapped fields). ACTIVECAMPAIGN – admin-editable: unlimited named scores, every rule's conditions (full segment builder), point value (click the inline "Add 10 points" text to change it), sign (add or subtract, e.g. subtract for countries you don't ship to), and per-rule points expiration (click inline "Never" → set a period; docs suggest 1-3 months for engagement; expiration is an all-or-nothing drop-off after the period, NOT gradual decay). Score thresholds live wherever they're consumed (segment conditions, "Score Changes" trigger), not in the score object. Fixed: no default scores exist (blank-slate; recipes marketplace fills the gap), scores must be toggled "Active" before automations can use them, deal scoring requires the Pipelines or Sales Engagement add-on on Plus/Professional/Enterprise plans (contact scoring is on all plans). Engagement Tagging recipe intervals are pre-filled but editable; it is forward-only ("cannot determine retroactive engagement").

CONFIG UX

ZOHO: wizard + table pattern, admin-only (Setup > Automation > Scoring Rules → "New Scoring Rule" button → popup: name, Type=Manual Score, module, layout, description → "Next" → criteria screen). Criteria are structured rows: field dropdown, operator dropdown, value, Add/Subtract, points number.

Touchpoint scoring is a checklist of channels, each expanding to its own actions with a points box per action. Ends with an opt-in prompt to map the six computed scores onto record fields. Result surfaces on the record detail page as a "Scoring Rules" related-list-style section with graphical representation (the old single-rule version showed one score pinned at the top of the record – the redesign traded glanceability for multi-rule capacity). It's a classic admin-console: powerful, but a settings-area artifact a non-technical event producer would need a consultant for (an entire Zoho partner ecosystem – Zenatta, Marks Group – exists to configure exactly this). ACTIVECAMPAIGN: sentence-style inline editing, much lower intimidation. Contacts > Scoring → "Add New Score" → modal picks Contact vs Deal score → name + description → "Add New Rule" opens the SAME segment builder used for segments/automations (leftmost dropdown = condition category, "Add another condition" stacks conditions) → after Save, the rule renders as clickable prose: the text "Add 10 points" edits points, the text "Never" edits expiration → "Active" toggle publishes. Per-event accrual is deliberately pushed into the automation canvas (visual flow builder: trigger → Runs: Multiple Times → "Adjust score" action modal with dropdowns for which score / add-reduce-set / points / Expires). Consumption UX: deals list "Add Filter" > "Deal Details" > pick score; per-stage sorting via stage gear icon > "Edit Stage" > "Deal Sorting" dropdowns. Non-technical admins are onboarded via pre-built recipes ("Create an automation" modal + Marketplace) with an "Automation Setup Wizard" that walks each prompt – the recipe, not a blank rule form, is the actual onboarding surface.

TAKEAWAYS FOR FLOWZ

- Ship opinionated defaults, not blank rule builders. Both platforms make admins start from zero and both compensate externally (Zoho: partner/consultant ecosystem; AC: recipe marketplace with pre-filled, editable intervals). Flowz should ship its lifecycle stages and tiers pre-configured with sensible event-world thresholds and let producers EDIT them – the AC Engagement Tagging pattern ('intervals included for you, you can change them') is the exact model, and it maps 1:1 onto New/Rising/Core/At-risk/Lapsed.
- One condition builder everywhere. AC's biggest UX win: scoring rules, segments, automations, and conditional content all use the identical segment builder, so learning it once unlocks everything. If Flowz adds editable thresholds for tiers, stages, AND tags, they should share one threshold/condition editor component, not three bespoke config screens.
- Keep score $\neq$ tier $\neq$ stage as separate constructs – both platforms do. Zoho separates attribute scores from behavioral touchpoint scores (six categories); AC separates Contact Score (who they are/engagement) from Deal Score (opportunity value). Flowz's two-dimension model (lifecycle = recency behavior, tier = cumulative value) is already the industry-correct decomposition; if a numeric score is added, make it a THIRD derived signal (engagement score feeding the lifecycle stage), not a replacement for tiers.
- Config that reads as a sentence beats config that reads as a form. AC's inline-editable prose ('Add 10 points' click-to-edit, expiration shown as the word 'Never') is far less intimidating for a non-technical producer than Zoho's field/operator/value/points grid. For Flowz's B2B clients, render tier thresholds as editable sentences: 'Gold = attended [5]+ events AND spent [2,000]+'.
- Model decay explicitly via negative points and point expiration – the two mechanisms both platforms use for at-risk detection. AC's expiration is a per-rule all-or-nothing drop-off after N days (docs suggest 1-3 months), not a gradual curve; Zoho uses explicit negative scoring (unsubscribe, opt-out at -50/-100, inactivity windows like 'in the last 20 days'). Flowz's At-risk/Lapsed transitions should be shown as an explicit, editable rule ('points from an event expire after 18 months') rather than a hidden decay function – legibility is the feature.
- Tags stay flat but need conventions + a manager. AC keeps tags flat and instead documents prefix conventions ('[ACTION] ...', 'Interest - Product - ...') and provides a Tag Manager showing per-tag contact counts, which automations use each tag, descriptions, and merge/prune tools. Flowz's EarlyBuyer/GroupBuyer/Spender/Responsive tags should get category prefixes (Behavior:/Value:/Channel:) plus a manager view with usage counts – and guardrails matter: Zoho's caps (10 rules per layout, 5 active) exist because unbounded config sprawls; give producers a bounded, curated editing surface, not an open rule engine.

D5 Braze · Iterable · CleverTap

Braze + Iterable + CleverTap (lifecycle-marketing engagement scoring survey)

ARCHITECTURE

Three distinct patterns. (1) BRAZE – no built-in lifecycle stages at all: everything is dynamic Segments composed from filters. Engagement signals are first-class filter primitives: "Last Engaged With Message" (last click/open on any channel), "Last Received Any Message", "Last Received Email/Push/SMS/WhatsApp", "Clicked/Opened Campaign", "Clicked/Opened Step" (Canvas), "First Used App"/"Last Used App", "Median Session Duration", session count. On top sits Predictive Churn (Predictive Suite): an ML model writes a Churn Risk Score 0-100 onto each user, bucketed into exactly 3 named bands – Low Risk 0-50, Medium Risk 50-75, High Risk 75-100 – and both "Churn Score" and "Churn Category" become segment filters like any other. Model trains on custom events, purchase events, eCommerce events, campaign interaction events, and session data; retrain automatically every 2 weeks. (2) ITERABLE – Brand Affinity: one AI-computed scalar ("propensity for future engagement") per contact, translated into 5 labels (Loyal/Positive/Neutral/Negative/Unscored) exposed as a contact property usable in segmentation, journeys, and dynamic content. Signals: opens and clicks across email, push, and in-app; crucially it scores the engagement RATIO (percent of received messages interacted with), not raw counts, with exponential time decay on older interactions, and the lookback window auto-adjusts to the brand's send cadence. Transactional-message engagement counts as a positive signal. Updated continuously. (3) CLEVERTAP – RFM Analysis: users are percentile-ranked on Recency and Frequency of one chosen event (scores 1-5 each), and the R×F grid maps to 10 named segments; Monetary is overlaid as Average Monetary Value per segment plus channel reachability. An "RFM Transition" chart shows user flow between segments over time (e.g., New Users → Champions). Separately, Intent Based Segmentation auto-creates 3 predictive micro-segments per goal: most likely / moderately likely / least likely (to purchase, churn, or uninstall). Structural takeaway: lifecycle stage and value tier are NOT separate verticals in these platforms – CleverTap fuses them into one 2-axis grid; Braze and Iterable reduce everything to one score with 3-5 named bands that feed the segment builder.

VERBATIM NAMING

- CleverTap RFM (10 segments, verbatim from docs.clevertap.com/docs/rfm): Champions [R4-5, F4-5]; Loyal Users [R3-4, F4-5]; Potential Loyalists [R4-5, F2-3]; New Users [R4-5, F0-1]; Promising [R3-4, F0-1]; Needing Attention [R2-3, F2-3]; About to Sleep [R2-3, F1-2]; At Risk [R1-2, F3-4]; Cannot Lose Them [R1-2, F4-5]; Hibernating [R1-2, F1-2]
- Iterable Brand Affinity (5 labels): Loyal; Positive; Neutral; Negative; Unscored (insufficient message history)
- Braze Predictive Churn bands: Low Risk (score 0-50); Medium Risk (50-75); High Risk (75-100) – plus raw 'Churn Score' 0-100 and 'Churn Category' as filter fields
- CleverTap Intent Based Segmentation: 'most likely' / 'moderately likely' / 'least likely' (to purchase / churn / uninstall)
- Braze engagement filters (verbatim): 'Last Engaged With Message', 'Last Received Any Message', 'Clicked/Opened Campaign', 'First Used App', 'Last Used App', 'Median Session Duration'

WHAT CUSTOMERS CAN EDIT

BRAZE (most customer-editable): the customer defines what churn MEANS in a filter builder with AND/OR logic – action type "do" or "do not" + event + time frame up to a 60-day window (example from docs: "do not start a session" within 7 days); picks the Prediction Audience (default "All Users", max 100M); picks update cadence (weekly/monthly standard, daily costs extra); max 5 concurrent predictions per workspace; minimum churned/non-churned data thresholds are shown during setup and block building if unmet. The 0-50/50-75/75-100 band boundaries themselves are FIXED, but the slider lets you target any arbitrary score range, so bands are advisory. CLEVERTAP: customer picks the anchor event (purchase, session, video played...), the date range (recommended under 512 days), and the monetary property (any numeric – money, minutes, videos watched); scoring is percentile-based so thresholds self-calibrate to each app's activity distribution – the 10 segment names and the R/F grid mapping are FIXED, no renaming or re-bucketing. ITERABLE: essentially zero customer tuning – label boundaries are "determined dynamically for each project" by Iterable's model; customers cannot edit thresholds, weights, or label names; Iterable monitors and adjusts the model itself. Pattern: platforms expose the DEFINITION (which event, which window, which audience) for editing, never the scoring internals or band names.

CONFIG UX

BRAZE: setup is a short wizard (define churn via filter rows + AND/OR, pick audience, pick cadence), then an analytics page does the heavy lifting – a histogram of the score distribution in 20 equal buckets with an interactive slider underneath; dragging the slider shows live audience counts and estimated churners/non-churners in the selected range, with "Create Segment" / "Create Campaign" buttons that materialize the selection as a reusable segment. Explainability comes from a Churn Correlation Table (which attributes/behaviors make users more/less likely to churn) and a Prediction Quality meter (lift-based, 60-100% labeled "excellent", 20-40% "fair"). All self-serve, admin-level, no consultant. CLEVERTAP: RFM is a full dashboard page – pick event + date range + monetary property, get the colored grid of 10 named cells showing user count, average monetary value, and channel reachability per cell; click a cell to spin off a segment or campaign; RFM Transition visualizes movement between named segments across periods. ITERABLE: no config surface at all – labels simply appear as a contact property on profiles and as a queryable field in the segmentation builder, Journey Studio branching, and dynamic-content templates. Three UX tiers: full rule-builder (Braze), parameterized analysis (CleverTap), invisible/automatic (Iterable).

TAKEAWAYS FOR FLOWZ

- Merge Flowz's lifecycle stages and tiers into one 2-axis construct rather than two parallel verticals: CleverTap proves a Recency x Frequency grid with named cells is instantly legible, and its 10 names map almost 1:1 onto Flowz's model (New Users~New, Potential Loyalists~Rising, Champions/Loyal Users~Core, At Risk/About to Sleep~At-risk, Hibernating~Lapsed) – keeping LTV/monetary as an overlay per cell, not a third axis.
- Use percentile-based scoring, not absolute thresholds, as the default: CleverTap ranks users 1-5 by percentile so a yoga-retreat brand with 2 events/year and a party brand with 50 both get sensible bands with zero configuration; offer absolute-threshold overrides as the advanced edit, not the starting point.
- Keep the numeric score internal, expose 3-5 named bands: every platform converts its score to plain-English labels (Low/Medium/High Risk; Loyal/Positive/Neutral/Negative) and B2B users act on labels – and always include an explicit 'Unscored' state for attendees with too little history (Iterable's fifth label) instead of silently defaulting them to the bottom band.
- Let clients edit the DEFINITION, never the model: Braze's churn builder (event + do/do-not + N-day window + audience filter) is the right editable surface for Flowz – 'Lapsed = no ticket purchase within 400 days' as an editable rule row – while band boundaries and scoring math stay fixed or auto-calibrated (Iterable hides them entirely on purpose).
- Score engagement RATIO with time decay, not raw counts: Iterable scores the percent of messages a contact interacted with and decays old activity exponentially, with the lookback auto-tuned to send cadence – for Flowz, score share of a producer's events attended / emails opened relative to opportunity, so attendees of low-frequency producers aren't unfairly down-ranked.
- The threshold-editing UX to copy is Braze's histogram + slider: show the client the live distribution of their attendees along the score, let them drag band boundaries and see counts change in real time, then one-click materialize a band as a segment ('Create Segment' button) – this makes custom thresholds self-explanatory instead of a blind numeric input.

D6 Event-Industry Platforms

Event-industry fan segmentation survey: Audience Republic, Tixr, vivenu, Ticket Fairy, DICE, Insomniac Passport, Live Nation (Festival Passport / Lawn Pass)

ARCHITECTURE

The dominant pattern across event platforms is NOT lifecycle-stages-plus-tiers. It is: (a) dynamic rule-based SEGMENTS built from behavioral filters, (b) manual TAGS for human labels, (c) per-fan raw metrics on the profile (LTV, total spend, events attended, email/SMS engagement), and (d) fan-facing loyalty TIERS/POINTS kept as a separate product, not a CRM construct. Per platform: AUDIENCE REPUBLIC – two constructs only: Segments ("dynamic groups based on filters like location, email engagement, or purchase history" that "update automatically as new fans match your criteria") and Tags ("manual labels you apply to contacts – great for categorising contest entrants, VIPs, or internal audience groups"). No numeric fan score; the fan profile surfaces lifetime value, purchase history, and engagement instead. Segments auto-drop fans when conditions change (e.g., after they buy). Superfans are just a saved segment you build (e.g., top spenders), then used for lookalike audiences on Meta/Google/TikTok. TIXR – raw data model: 90+ data points per fan (purchase history, attendance, demographics, preferences) in its "Studio" backend; segmentation itself happens by filters (ticket tier, event, venue, artist, genre, purchase history) or downstream in an integrated CRM/CDP. The only explicit scoring is Tixr Rewards: referral POINTS toward organizer-defined reward tiers, with optional leaderboards – points-as-score, earned by actions, not a computed 0-100. VIVENU – launched "Customer Segments" (Aug 2025): dynamic segments built with "Custom Rules" over spend, tickets purchased, attendance, promo usage, location, donations, email engagement, with AND/OR logic; "Dynamic segments work as tags across your customer journey"; each segment shows pre-analyzed metrics (total value, average basket size, "spending tiers", top locations). No RFM or scoring model exposed. Paired with "vivenu Engage" for campaigns and automated fan rewards. DICE – thin publicly: fans "follow" venues/promoters in the app, promoters get "audience insights" and analytics, 40%+ of sales come from algorithmic discovery; no public segment taxonomy. SEE TICKETS – no meaningful public CRM segmentation docs found; move on. LOYALTY PROGRAMS – Insomniac Passport is an invite-only paid membership (Insomniac picks members from its data; criteria opaque, anecdotally random/attendance-based) with three access tiers priced \$50/mo (California festivals), \$60/mo (US incl. EDC), \$80/mo (US Max – everything), plus \$20/event and up to 3 guests, priority/dedicated entrance, 10% merch discount. Live Nation runs product-passes, not computed tiers: Festival Passport (\$999, 100+ festivals, added a VIP tier), Lawn Pass/"Lawnie Pass" (\$99-\$399 per amphitheater; \$239 at Blossom 2024, guaranteed lawn + Fast Lane), Club Pass.

VERBATIM NAMING

- Audience Republic: "Segments" and "Tags" (the only two constructs); template segments: "Past Buyers (Not This Year)", "Presale Registrants Who Haven't Purchased", "Top Spending Customers", "Fans in a Specific Location"; marketing language: "superfans", "top ticket buyers", "most loyal fans"
- Tixr: "Tixr Rewards" points, organizer-named reward tiers, "leaderboards", "superfans" (in case studies – e.g., High Sierra 2025 superfans referring 20+ tickets each)
- vivenu: "Customer Segments", "Custom Rules", "Dynamic & Real-Time Segments", "spending tiers", segment types cited: "loyal season ticket holders", "one-time VIP buyers", "fan club & membership groups"
- Ticket Fairy (blog, reflects industry practice): "Super fans" (attend nearly every week), "Frequent attendees" (3+ events in 12 months), "First-timers", "Lapsed attendees" (no return in 12 months), "High spenders" (VIP upgrades/premium seats), "Locals vs. out-of-towners", "Gold Fan" (10+ events milestone), "Ambassadors" (referrers)
- Insomniac: "Insomniac Passport" with tiers "California" / "US" / "US Max"
- Live Nation: "Festival Passport" (+ "VIP tier"), "Lawn Pass" / "Lawnie Pass", "Club Pass"
- DICE: "followers", "audience insights", "Groups" (friends attending together – the group-buyer construct)

WHAT CUSTOMERS CAN EDIT

Everything behavioral is customer-editable; nothing is a fixed taxonomy. AUDIENCE REPUBLIC: customers build any segment from any filter combination with layered AND/OR logic (e.g., criterion "Ticket Name is 'Early Bird'" = early-buyer segment), name it, save it; edit/delete via a three-dot menu. Tags are freeform text created inline. No fixed lifecycle stages or tiers exist to override. TIXR: organizers configure which ticket types earn points, the point value per ticket tier, the point thresholds that unlock each reward tier, and the reward type per tier (cash-back, merch, upgrades, experiences); leaderboards are optional. TVIVENU: customers define segments via Custom Rules over any event/purchase data; segments auto-update ("no spreadsheets or exports needed") and double as tags; segment-level metrics are computed for you (not configurable). Thresholds seen in industry practice (Ticket Fairy): 3+ events/12mo = frequent, 4 shows = loyalty reward trigger, 8+ events/yr = VIP status, 10+ events = "Gold Fan", 12 months inactive = lapsed, ~50% repeat-attendance baseline. What is NOT customizable anywhere: no platform exposes editable scoring weights or a tunable composite fan score – the closest is Tixr's points-per-action config.

CONFIG UX

Universal pattern: a filter/rule builder, not weight sliders and not consultant-configured. Audience Republic: Audience > Filters > add criteria > "Save Segment" with a clear name; segments live in Audience Manager next to tags; tags applied by selecting contacts > Tag icon > Create New. vivenu: in-platform builder framed as "ask any question of your audience data, and instantly see the answer as dynamic groups", with instant per-segment analytics (size, total value, avg basket, revenue share) shown as you build – the comprehension trick is showing the money metrics next to the rule. Tixr: rewards configuration screen where the organizer sets point values and tier thresholds directly; fulfillment split into Tixr-fulfilled (auto cash-back) vs client-fulfilled (merch/experiences). All platforms make segments dynamic/self-updating by default – a fan falling out of criteria leaves the segment automatically. Nothing is admin-only or professional-services-gated in the products reviewed.

TAKEAWAYS FOR FLOWZ

- Flowz's 3-vertical model (lifecycle + tiers + tags) is already richer than what competitors ship – but competitors win on editability. Industry standard is: tags stay manual (Flowz matches), while everything behavioral is a saved, named, dynamic rule (filters + AND/OR) the client can open and edit. Flowz's lifecycle stages and tiers should each be expressible as visible, editable threshold rules, not black boxes.
- Use plain-English descriptive names for internal segments – the industry vocabulary B2B clients already know is: First-timers, Frequent attendees, Super fans, High spenders, Lapsed. These are instantly legible; save metal names (Bronze/Gold) or branded names (Passport, Gold Fan) for FAN-FACING loyalty tiers, where every real program (Insomniac Passport, Festival Passport, Tixr Rewards tiers) uses branded/aspirational names the producer defines.
- Default thresholds should match industry convention and be small integers of events + recency windows: 3+ events/12mo = frequent/core, 8-10+ = top tier/VIP, 12 months inactive = lapsed, first purchase = new. Ship these as editable defaults – Ticket Fairy's published numbers are the closest thing to an industry baseline.
- Skip an abstract 0-100 fan score. No event platform exposes one; they surface raw LTV + events attended + engagement on the fan profile and let clients rank/filter on them. The only 'score' clients understand instantly is earned points for actions (Tixr Rewards: organizer sets points per ticket type and points-per-reward-tier thresholds) – if Flowz adds scoring, make it action-based points feeding tiers, with client-editable point values.
- EarlyBuyer/GroupBuyer/Spender-style tags are validated, but competitors derive them from behavior rules rather than static labels: early buyer = 'Ticket Name is Early Bird' or bought in first sale window; group buyer = multi-ticket orders or DICE-style Groups; advocate = has referral conversions. Make Flowz's behavioral tags auto-computed (dynamic) and keep manual tags as a separate freeform layer, mirroring Audience Republic's Segments-vs-Tags split.
- The comprehension pattern that makes segments 'instantly understandable' to B2B clients is vivenu's: show live metrics (member count, total spend, revenue share, avg basket) directly beside each segment definition as the client edits the rule – the segment explains itself in money terms. Pair that with auto-updating membership (fans drop out when they stop qualifying), which every platform treats as table stakes.

D7 The RFM(E) Model

Canonical RFM(E) segmentation model – synthesized from practitioner implementations: Putler (canonical 11-segment table), CleverTap (scoring mechanics + weighted/RFM-E), Omniconvert (alternative naming + scale sizing), Bloomreach Engagement (auto-quantile product), WebEngage (configurable RFM grid), Spektrix/ArtsManaged (arts-ticketing seasonal adaptation)

ARCHITECTURE

RFM is a SCORE-DRIVES-SEGMENT architecture, not independent verticals. Layer 1: three numeric scores 1-5 per customer – Recency (days since last purchase), Frequency (purchase count), Monetary (total/avg spend) – assigned by quintiles over the actual customer distribution (top 20% of each dimension = 5, next 20% = 4, etc.; CleverTap). This yields up to 125 combined codes from 111 to 555. Layer 2: the 125 codes collapse into ~11 named segments via range rules on a 2D grid of Recency x (Frequency+Monetary averaged). The canonical Putler mapping: Champions R4-5/FM4-5; Loyal Customers R2-5/FM3-5; Potential Loyalist R3-5/FM1-3; Recent (New) Customers R4-5/FM0-1; Promising R3-4/FM0-1; Customers Needing Attention R2-3/FM2-3; About To Sleep R2-3/FM0-2; At Risk R0-2/FM2-5; Can't Lose Them R0-1/FM4-5; Hibernating R1-2/FM1-2; Lost R0-2/FM0-2. Segments are mutually exclusive, derived, and recomputed on a schedule (Bloomreach: monthly repeat-trigger by default; quantile cut-points drift with the base automatically – no manual recalibration). Note the model deliberately fuses lifecycle AND value into ONE grid: 'At Risk' vs 'Can't Lose Them' is the same recency band split by value – exactly the lifecycle-x-tier cross Flowz currently models as two separate verticals. Extensions: (a) RFM-E / RFE – CleverTap: Monetary is replaced (or supplemented) by an Engagement factor (visit duration, sessions, email/push response) for indirectly-monetized or low-purchase-count businesses; (b) Weighted RFM – instead of equal thirds, weight dimensions per industry: common practice 50% Recency / 30% Frequency / 20% Monetary for repeat-purchase retail; gaming weights Frequency up; subscriptions weight Recency around renewals. (c) Scale sizing – Omniconvert: use 1-3 bins under ~30k customers, 1-4 for 30k-200k, 1-5 above 200k (with 1-4 events/year and small B2B-client bases, quintiles over-slice; tertiles or quartiles like MCP Analytics' NTILE(4) are more honest). Event/ticketing adaptation of Frequency: purchase cadence is 1-4/year, so (1) redefine the unit – 'events attended per year' or 'seasons attended' instead of raw orders (ArtsManaged); (2) extend the analysis window to 18-24 months so customers get 4-6 purchase opportunities before scoring frequency (MCP Analytics); (3) don't demote seasonal customers mid-offseason – a buyer who only buys each festival season is Core, not At-risk (marketing.link); (4) quantile-based scoring absorbs seasonality automatically because cut-points come from the population, not calendar rules; (5) Spektrix's arts benchmark: 'Frequent Attender' = 3+ visits in the last 12 months, 'Recent' = activity in last 6 months, segments mutually exclusive by priority order.

VERBATIM NAMING

- Canonical 11 (Putler): Champions
- Loyal Customers
- Potential Loyalist
- Recent Customers (a.k.a. New Customers)
- Promising
- Customers Needing Attention
- About To Sleep
- At Risk
- Can't Lose Them
- Hibernating
- Lost
- Bloomreach variant: Lost customers / Hibernating customers / Cannot lose them / At risk / About to sleep / Need attention / Promising / New customers / Potential loyalists / Loyal / Champions
- WebEngage variant (ascending order): Lost, Hibernating, About to Sleep, At Risk, Cannot Lose Them, Need Attention, New Users, Promising, Potential Loyalist, Loyal, Champions
- Omniconvert relationship-metaphor set: Soulmates, Lovers, Potential Lovers, New Passions, Flirting, Apprentice, Platonic Friends, About To Dump You, Ex-Lovers, Don Juan, Break-Ups
- Spektrix arts 4-segment: Frequent Attenders, Recent Attenders, Locals, Lapsed Attenders & Everyone Else

WHAT CUSTOMERS CAN EDIT

Across major platforms, LESS is editable than you'd expect. Bloomreach: ready-to-use 'black box' – quantile slicing is automatic, segment names fixed, bin boundaries NOT user-editable; only the refresh cadence (repeat-trigger node, default monthly) is configurable; docs explicitly warn 'Dynamic RFM is editable but due to its design, it is not easy to make changes' (requires editing underlying Jinja logic, not UI).

WebEngage: user configures the INPUTS, not the math – pick the time frame (default 30 days), pick which event counts as Recency/Frequency (any custom or system event, e.g. 'Session Started', with filters), pick the Monetary event + attribute (e.g. 'Checkout Completed' aggregated by 'Cart Value'; monetary optional – becomes RFE); the 1-5 scoring ranges and segment thresholds are predefined and non-editable. Putler: fully automatic 11 segments; user can recalculate over custom date ranges but not edit thresholds or names.

CleverTap: automated grid + segment-transition view; simplifies to a 5x5 (25-cell) Recency-x-Frequency grid rather than exposing 125 codes. Pattern: platforms fix the SEGMENT LOGIC and names (that's the product's opinion/value) and let customers edit the DEFINITIONS feeding it – which events count, over what window, what money field. Custom/freeform segmentation is offered as a separate parallel feature (Putler 'Custom Segmentation', Spektrix saved criteria sets), never by mutating the RFM grid.

CONFIG UX

Editing UX patterns observed: (1) Query-bar config (WebEngage) – a top bar with three dropdowns: time frame, R/F event picker, monetary event+attribute picker; hit apply, grid recomputes; segment tiles sized proportionally to user counts; empty segments simply don't render. (2) Grid visualization as the primary UI (CleverTap, WebEngage) – a 2D heatmap (Recency vs Frequency) where each named segment is a colored region; per-cell counts, reachability, and avg monetary value; plus a 'Transition' view showing users flowing between segments over time – the movement view is itself a selling point to B2B clients. (3)

Scenario/recipe-based (Bloomreach) – RFM ships as a prebuilt scenario the customer instantiates; power users can fork the underlying logic but the vendor discourages it. (4) No platform found exposes weight sliders or per-segment threshold editors in mainstream UI – weighting is a consultant/data-team pattern (done in SQL/notebooks), not an end-user control. (5) Arts-ticketing practice (Spektrix) is consultant-assisted: plain-English segment definitions ('3+ visits in last year'), enforced mutual exclusivity via priority order, deliberately only 4 segments because venue marketers can't action 11.

TAKEAWAYS FOR FLOWZ

- Merge Flowz's lifecycle and tier verticals into one derived grid: canonical RFM already encodes lifecycle (Recency axis) x value tier (F+M axis) in a single model – 'At Risk' vs 'Can't Lose Them' is literally the same lapsing-recency band split by spend. Keep Bronze/Silver/Gold as a display layer over the FM score if producers like medals, but derive both from the same scores so they can never contradict.
- Use behavior-descriptive segment names, not metal tiers: every major platform converged on self-explanatory names (Champions, At Risk, Can't Lose Them, About To Sleep) because a B2B client understands 'At Risk' instantly but must be taught what 'Silver' means. Omniconvert proves naming is skinnable (Soulmates/Ex-Lovers) – Flowz could offer event-flavored names (Headliners, Regulars, First-Timers, Ghosting, No-Shows) as a skin over canonical logic.
- Score with quantiles over the producer's own attendee base, not fixed thresholds: quintile cut-points auto-adapt to each client's size and seasonality with zero manual recalibration (Bloomreach/MCP Analytics). But with 1-4 events/year, use 3-4 bins, not 5 (Omniconvert's rule: 1-3 bins under ~30k customers) – a 5-bin Frequency over 'attended 1, 2, or 3 events' is fake precision.
- Adapt the model for event cadence explicitly: Frequency = events attended per rolling 24 months (gives 4-6 purchase opportunities per MCP Analytics); Recency measured in 'events missed since last attendance' or seasons, not days; never auto-demote between seasons (a buyer of every annual festival is Core in month 11). Spektrix's arts benchmark: 3+ visits/12 months = frequent, 6 months = recent.
- Let clients edit INPUTS and labels, not the math: the winning config UX (WebEngage) is three dropdowns – time window, which events count, which money field – while scoring bins and segment logic stay fixed. No mainstream platform ships weight sliders or per-segment threshold editors; Bloomreach explicitly warns against editing the logic. For Flowz: expose window + event-type filters + renameable labels + editable tier boundaries at most; keep the segment mapping as the product's opinion.
- Add Engagement as the 4th dimension (RFM-E) using Flowz's existing tags: CleverTap's RFM-E swaps/augments Monetary with engagement (opens, responses, app sessions) – Flowz's 'Responsive' tag is already an E signal. Use E to split identical RF scores (e.g. lapsed-but-still-opening-emails = 'At Risk, reachable' vs lapsed-and-silent = 'Lost'), and keep tags (EarlyBuyer, GroupBuyer) as orthogonal descriptors, never as segment inputs – no platform mixes freeform tags into the RFM grid.

D8 Loyalty-Tier Naming

Cross-brand loyalty tier naming research (Sephora, Marriott Bonvoy, Delta Medallion, Starbucks, Amex, oneworld/Star Alliance, lululemon, Ulta, e.l.f., Nike, Apex Legends ranked ladder)

ARCHITECTURE

Every major program converges on the same skeleton: ONE ordered status dimension, kept strictly separate from the spendable points currency (Sephora points vs Insider/VIB/Rouge status; Starbucks Stars vs Green/Gold/Reserve level; Delta miles vs Medallion status). Base tier is free and universal; 2-5 earned tiers sit above it; each tier boundary is 1-2 published numbers over an annual qualification window (Sephora: \$350/\$1,000 calendar-year spend; Marriott: 10/25/50/75/100 nights + \$23K for Ambassador; Delta: \$5K/\$10K/\$15K/\$28K MQD; Starbucks: 500/2,500 Stars; Ulta: \$500/\$1,200; e.l.f.: 1,000/2,000 points). Status requalifies annually – the 'fear of falling' is a deliberate retention mechanic. Five naming PATTERNS emerged: (1) METAL LADDER (Marriott Silver→Titanium, Delta Silver→Diamond, Amex Green→Platinum, Apex Bronze→Diamond) – culturally pre-ranked, zero learning cost, but generic; (2) GEM LADDER (oneworld Ruby/Sapphire/Emerald) – elegant but FAILS the stranger test: nobody knows Emerald outranks Sapphire without a legend; (3) INSIDER-LANGUAGE LADDER (Sephora Insider→VIB→Rouge; e.l.f. Fan→Pro→Icon; generic Member→Insider→VIP) – exclusivity vocabulary itself encodes rank; (4) BRAND-FLAVORED/JOURNEY LADDER (lululemon Collective→Plus→Pinnacle; Starbucks Green→Gold→Reserve mixing brand color + metal + exclusivity; status-verb ladders like Launch/Grow/Scale from membership-site practice); (5) APEX-BREAK: the top tier deliberately breaks the naming scheme to signal 'beyond the ladder' – Rouge, Centurion, Ambassador, Reserve, Apex Predator (which isn't even a threshold: it's the live top 750 Masters per platform). Psychology research (Drèze & Nunes, 'Feeling Superior', Journal of Consumer Research 2009): enlarging the top tier DILUTES perceived status; adding a subordinate bottom tier ENHANCES the status of the tier above it; non-qualifiers prefer multi-tier hierarchies; gold/silver labels alone signal a selective hierarchy. Industry data: tiered programs drive ~27% higher annual spend and 37% of customers spend more specifically to reach a status level.

VERBATIM NAMING

- Sephora Beauty Insider: Insider / VIB / Rouge (\$0 / \$350 / \$1,000 annual spend)
- Marriott Bonvoy: Member / Silver Elite / Gold Elite / Platinum Elite / Titanium Elite / Ambassador Elite (0/10/25/50/75/100 nights, Ambassador also \$23K spend)
- Delta SkyMiles: Silver / Gold / Platinum / Diamond Medallion (\$5K/\$10K/\$15K/\$28K MQD, thresholds held flat 2025-2027)
- Starbucks Rewards (2026 relaunch): Green / Gold / Reserve (0 / 500 / 2,500 Stars)
- American Express: Green / Gold / Platinum / Centurion (invite-only 'Black Card' apex)
- oneworld: Ruby / Sapphire / Emerald; Star Alliance: Silver / Gold
- Apex Legends ranked: Rookie / Bronze / Silver / Gold / Platinum / Diamond / Master / Apex Predator (Predator = live top 750, not a threshold)
- lululemon Membership: Collective / Plus / Pinnacle
- Ulta Ultimate Rewards: Member / Platinum / Diamond (\$0 / \$500 / \$1,200)
- e.l.f. Beauty Squad: Fan / Pro / Icon (0 / 1,000 / 2,000 points)
- Nike Membership: single free 'Member' tier – deliberate no-ladder counterexample

WHAT CUSTOMERS CAN EDIT

In consumer programs the member edits nothing – but the OPERATOR pattern is the lesson: names are the durable brand layer, numbers are the tunable layer. Delta held Medallion names constant while repricing MQD thresholds repeatedly; Sephora has kept Insider/VIB/Rouge stable since 2013 while adjusting perks; Starbucks re-tiered its entire program in March 2026 (flat → Green/Gold/Reserve) without touching its points currency. Membership-industry naming data (Communal's 500+ org survey) shows descriptive names ('Family', 'Regular') vastly outnumber branded ones (Gold/Silver/VIP appear <25 times combined) in small orgs – small B2B clients default to descriptive because their staff must decode the ladder instantly. For Flowz: tier NAMES should be per-client editable text (skins) over fixed internal keys (tier_1..tier_5); tier THRESHOLDS should be editable numbers on the two existing axes (events attended + LTV) with an annual/rolling window setting; the NUMBER of tiers should be 3-5 with a locked ordering, not free-form.

CONFIG UX

The implied config surface, from how these programs publish rules: an ordered tier table – each row = name + color/icon + 1-2 numeric thresholds (min events, min LTV) + qualification window – where renaming is presentation-only and automations/flows reference the stable tier key, so a yoga client renaming tier_4 from 'Headliner' to 'Ambassador' breaks nothing. Threshold edits should show a live preview of how many attendees land in each tier (the Drèze & Nunes finding makes this actionable: warn when the top tier exceeds ~5-10% of the base, because an oversized top tier destroys its own value). Offer 3-4 prebuilt name skins (see ladders below) selectable at workspace setup, with metals as the fallback. Loyalty SaaS vendors (Open Loyalty, LoyaltyLion class) expose exactly this: rule-builder rows of 'spend X in window Y → tier Z' – admin-only, no end-customer visibility into mechanics.

TAKEAWAYS FOR FLOWZ

- Naming rules that survived the research: (a) a stranger must rank the ladder in under 2 seconds – metals pass, gems fail (Ruby/Sapphire/Emerald needs a legend); (b) 1-2 words per name; (c) let the TOP tier break the naming pattern to feel 'beyond the ladder' (Rouge, Centurion, Reserve, Ambassador, Apex Predator); (d) never share vocabulary between the lifecycle axis and the tier axis – Flowz's 'Core' lifecycle stage means no tier may ever be named Core; (e) keep the top tier small and keep a subordinate bottom tier – both measurably raise perceived status (Drèze & Nunes, JCR 2009).
- Architecture: ship ONE default ladder on stable internal keys (tier_1..tier_5); client-editable name skins + numeric thresholds (events attended, LTV, qualification window), exactly the Delta model of fixed names over repriced numbers. Add annual requalification as an option – status that can be lost is the strongest retention mechanic in the research (37% spend more to reach/keep status).
- Universal descriptive ladders (work for any client vertical, instantly readable by B2B staff): LADDER 1 'First-Timer / Returning / Regular / Superfan / Legend' – pure fan-intensity English, matches how promoters already talk; LADDER 2 'Member / Insider / VIP / Inner Circle' – Sephora-pattern exclusivity vocabulary encodes the order itself.
- Music-flavored skins (bass festivals, party brands): LADDER 3 'GA / Front Row / Backstage / Headliner' – venue proximity physically encodes rank, every promoter and attendee decodes it instantly; LADDER 4 'Spark / Glow / Blaze / Inferno' – fire-size is monotonic and on-brand for EDM/bass, deliberately wrong for yoga (that's what skins are for).
- Wellness skin + safe fallback: LADDER 5 'Explorer / Regular / Devotee / Ambassador' – retreat-friendly, with Marriott's service-flavored apex; LADDER 6 'Bronze / Silver / Gold / Platinum (+ Diamond)' – keep the current metals as the default skin since decode rate is 100%, and the research says Flowz's real problem isn't these names, it's that tier vs lifecycle vs tags are three unlabeled dimensions colliding in one UI.
- Dynamic-apex option worth building: LADDER 7 'Rising / Established / Elite / Apex' where Apex is defined as 'top N% (or top N) of this client's attendees per season' – the Apex Predator top-750 pattern. This gives a 300-person yoga retreat and a 50,000-person festival an equally exclusive top tier with zero threshold tuning, and it's a genuine differentiator no CRM competitor ships as a first-class construct.

D9 Failure Modes

Failure-mode research (cross-domain): B2B lead scoring (HubSpot/Marketo-style), SaaS customer health scores (Gainsight-style), RFM segmentation (Klaviyo/Omniconvert-style), and tiered loyalty programs (airline/hotel/retail) – practitioner criticism and academic studies

ARCHITECTURE

Six recurring structural failure modes. (1) **ARBITRARY WEIGHTS, NO VALIDATION**: lead-scoring points are guesses nobody backtests – "a whitepaper download is worth 15 points" and "100 points = sales-ready" only because someone picked a round number in a meeting; when teams finally pull closed-won data, the rewarded signal often has zero relationship to revenue (Pedowitz Group, Verum, ORM). ~68% of B2B companies use lead scoring but only ~40% of salespeople get value from it – the consuming team ignores scores built without their input and with no feedback loop. (2) **COMPOSITE SCORES HIDE THE SIGNAL**: a single blended health score has no drivers behind the number ("unhealthy, but never why") and no trajectory; documented case of an account showing green in the CSP that renewed at 78% of forecast (Querri). Teams keeping 4+ visible dimensions instead of one blend report ~34% better churn prediction (Accoil). (3) **RFM MISFIT FOR SEASONAL/LOW-FREQUENCY**: with an annual festival, recency/frequency scores collapse in the off-season – "customers swing segments due to seasonality, not intent," and scores change depending on when the analysis is run; RFM is explicitly called unsuitable for products bought only a few times, and long-cycle industries must recalibrate recency to realistic repurchase windows. This is Flowz's single biggest exposure: a 90/365-day 'Lapsed' clock is wrong for producers who run 1-2 events per year. (4) **THRESHOLD CLIFFS AND DEMOTION**: demotion from a tier has a negative, **ASYMMETRIC** effect on loyalty – demoted customers end up less loyal than never-promoted ones (Wagner/Hennig-Thurau/Rudolph, Journal of Marketing 2009); mitigations that measurably work: advance notice, short re-qualification windows, gain-framed messaging. Conversely the goal-gradient effect (Kivetz/Urmitsky/Zheng 2006) shows purchase acceleration near a visible threshold – cliffs cut both ways. (5) **TIER INFLATION**: when top tiers get crowded (Hilton Diamond via a credit card, Marriott Platinum via Amex Brilliant, AA Executive Platinum via card spend), finite perks dilute, then programs devalue (remove perks, raise thresholds), which enrages the exact customers the tiers were meant to retain. (6) **OVER-SEGMENTATION**: dozens of near-duplicate segments nobody actions; practitioner consensus lands on 4-7 (max 3-8) operating segments and ≤5 tiers; beyond that, personalization degrades to token level and ops time goes to list maintenance.

VERBATIM NAMING

- RFM auto-segment names in tools: Champions, Loyal Customers, Potential Loyalist, New Customers, Promising, Need Attention, About to Sleep, At Risk, Can't Lose Them, Hibernating, Lost
- Loyalty tiers: Bronze / Silver / Gold / Platinum (generic metals – instantly understood ordinal scale)
- Hotel/airline: Hilton Honors Diamond; Marriott Bonvoy Platinum Elite, Titanium Elite, Ambassador Elite; American Airlines Executive Platinum
- Lead scoring: MQL / SQL, HubSpot Score (0-100 with positive/negative attributes)
- Health scores: Red / Yellow / Green (traffic light), Gainsight scorecard measures
- Gamified relabels that failed: points as 'coins', tiers as 'levels' – documented as cosmetic gamification whose engagement decays within weeks (overjustification effect)

WHAT CUSTOMERS CAN EDIT

What real products let admins edit, and where it breaks: HubSpot manual lead scoring exposes free-form positive/negative criteria with arbitrary point values and no validation step – the canonical outcome is a score sales distrusts. Gainsight-style scorecards expose weight sliders that must sum to 100% – false precision, set once and never revisited. Loyalty platforms (Smile.io/LoyaltyLion pattern) expose tier thresholds as points-or-spend numbers; edits silently reshuffle the customer base with no preview of who moves. RFM tools mostly hard-code quintiles and segment names – customers can't rename 'Hibernating' even when it's wrong for their cadence. The failure common to all: threshold/weight edits ship without (a) a preview of population impact, (b) any backtest against outcomes, (c) an audit trail of what changed. For Flowz this means the editing surface should be RULES, not weights: 'Lapsed = missed [2] consecutive events they previously attended' is editable, explainable, and previewable; 'Engagement weight = 0.35' is none of those.

CONFIG UX

The config UX that survives contact with real B2B users: rule-builders with plain-language sentences and inline editable numbers, not weight matrices. Every segment assignment must answer 'why' on hover (the rule that fired, with the attendee's actual values vs. threshold). Threshold edits need an impact preview before save – 'this moves 1,240 attendees from Core to At-risk' – because silent reshuffles destroy trust exactly the way unvalidated lead scores do. Trend/direction matters more than the label: a static stage with no arrow (improving/declining) 'tells you almost nothing about when to act' (Querri criticism of health scores). Admin-only editing with sane locked defaults beats fully open config: practitioners consistently report that freely editable weights get set once in a kickoff meeting and never touched again, going stale within two quarters. If Flowz ever adds a numeric score, it must ship with its 2-3 top drivers displayed beside it – never the bare number.

TAKEAWAYS FOR FLOWZ

- Make recency EVENT-RELATIVE, not calendar-relative. 'At-risk/Lapsed' must be defined as missed editions ('skipped 2 consecutive events of a series they previously attended'), never 'no purchase in N days' – RFM's documented seasonal failure (customers swinging segments due to off-season, not intent) is fatal for annual festivals. Default the cadence per producer (detected from their event calendar) and expose the missed-events count as the editable threshold.
- Keep lifecycle and tier as two separate visible dimensions; never collapse them into one composite 'attendee score'. Single blended scores hide the actionable driver (the green-account-renews-at-78% case) and multi-dimension views predict churn ~34% better. If a score is added later, it must always render with its top 2-3 drivers and a trend arrow ('At-risk ↓: 2 missed events, opens down 60%'), and every stage badge needs a 'why' tooltip showing the fired rule.
- One segment = one default action, hard cap ~4-7 lifecycle stages and 4-5 tiers. If two segments would trigger the same campaign, merge them. Over-segmentation research is unanimous: beyond 3-8 operating segments, producers stop actioning any of them. Flowz's current 5 stages + 4 tiers is already at the healthy ceiling – resist adding more; add tags instead (tags are filters, not segments).
- Name tiers boring and ordinal, stages as plain behavioral states. Bronze/Silver/Gold/Platinum needs zero explanation – keep the metal scale or an equally self-ranking one; gamified relabels (coins, levels, VIP-vibes names) are documented to decay within weeks and confuse B2B users. For stages, prefer self-describing states a producer parses in one read (e.g., New / Returning / Regular / Cooling / Lapsed) over cute RFM names – 'Rising' and 'Core' need a legend; 'Returning' doesn't.
- Design tier transitions without cliffs. Demotion damages loyalty asymmetrically (demoted customers end up worse than never-promoted – J. Marketing 2009), so: require sustained qualification (events attended AND LTV over trailing window, so one VIP table can't buy Platinum), give advance notice + a short re-qualification grace window before demoting, frame demotion comms as gain ('attend 1 event to keep Gold'), and consider freezing tier during a producer's off-season. Show progress-to-next-tier to exploit the goal-gradient effect on the positive side. Guard against tier inflation with either a percentile-based top tier or a threshold review when top-tier share exceeds ~5-10%.
- Ship transparent rules with validated defaults, not editable weights. Lead scoring fails because weights are meeting-room guesses nobody backtests and the consuming team ignores the result (68% adoption, 40% perceived value). Flowz's config screen should be sentence-style rule builders ('Gold = [5]+ events AND [\$800]+ LTV') with two mandatory behaviors: an impact preview on every threshold edit ('moves 1,240 attendees Core→At-risk') and a periodic validation nudge comparing segment definitions against actual outcomes (e.g., '% of At-risk who actually lapsed'). If a rule can't be explained in one sentence and previewed, it doesn't ship.