

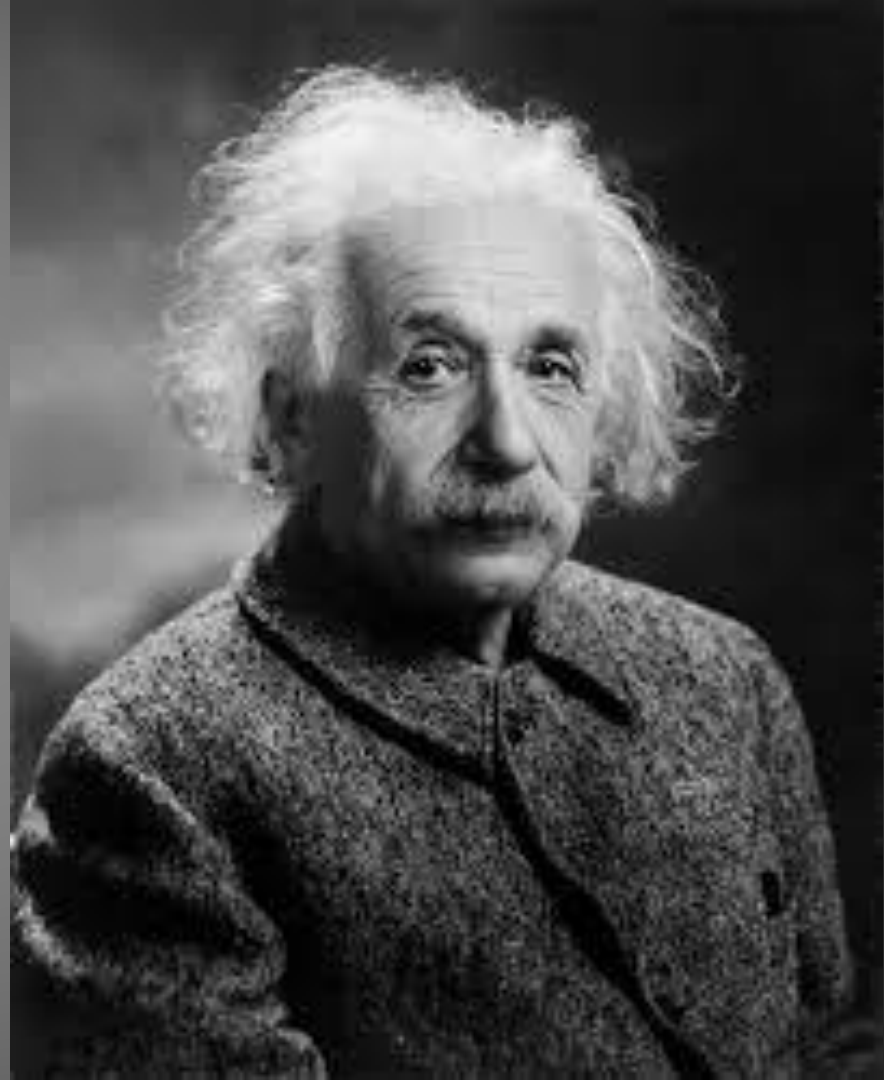
The Evolving Travel Ecosystem

New Assessments for New Platforms

Will Tate | Partner, GoldSpring Consulting

will@goldspringconsulting.com

“We cannot solve our problems with the same thinking we used when we created them.”



Why Traditional Evaluation Models Are Lacking

New Pricing Models

Transaction fees, subscription tiers, hybrid models — traditional cost comparisons no longer capture true program economics.

Fractured Content

NDC vs. EDIFACT vs. ATTRIBUTE BASED BUYING splits availability across channels. Legacy scorecards can't measure what travelers actually see at point of purchase.

UX as a KPI

Booking experience now directly drives online adoption rates. Traveler satisfaction must be scored, not assumed.

Future Risk

Roadmap credibility, financial stability, and innovation velocity matter as much as today's capabilities in a fast-moving market.

Why the Framework Matters

How the six-dimension model changes the conversation on both sides of the table

Why It Matters to Buyers

- Replaces a broken checklist with a weighted, evidence-based decision model
- Weighting the dimensions is a strategic lever that reflects program priorities
- Surfaces true P&L impact — servicing cost and adoption, not just headline fees
- Forces the questions RFPs avoid — native support and written roadmap commitments

Why It Matters to Sellers

- Changes what winning looks like — live performance beats polished demos
- Rewards transparency, genuine NDC / Attribute Based Buying maturity, and roadmap credibility
- Signals exactly which capabilities buyers will now measure
- Turns product and pricing investment into clear differentiation

The Scoring Methodology — How We Did It

01

Define Program Priorities

Establish what matters most to your travelers and program — weighting inputs, not vendor defaults.

02

Build the Scoring Rubric

Assign weighted criteria across all six dimensions. Scores are 1–10. ; CSAT items carry a 2× multiplier.

03

Gather Evidence

Demos, RFP responses, reference calls, and live data capture — not marketing decks.

04

Score Independently

Evaluators score blind to prevent anchoring. Results are aggregated and variance is discussed openly.

05

Validate & Calibrate

Weighted scores are cross-checked against real-world data (adoption, service metrics, cost per transaction).

The New Evaluation Framework

Six dimensions that go beyond the traditional RFP

1

**Content
& NDC /
Attribute
Based
Buying**

True availability
across all
channels

2

**Servicing
Impact**

Agent model,
touchless rate,
cost per tx

3

**UX &
Adoption**

Booking
experience
scored by
traveler

4

**Technology
Stack**

Integration
depth, API
maturity,
roadmap

5

**Commercial
Model**

Pricing
transparency,
incentive
alignment

6

**Innovation
& Risk**

Stability, AI
readiness,
future-proofing

Assessment by Framework Component

1

Content and NDC / Attribute Based Buying

Establish what matters most to your travelers and program — weighting inputs, not vendor defaults.

2

Servicing Impact

Assign weighted criteria across all six dimensions. Scores are 1–10; duty-of-care items carry a 2× multiplier.

3

UX and Adoption

Demos, RFP responses, reference calls, and live data capture — not marketing decks.

4

Technology Stack

Evaluators score independently to prevent anchoring. Results are aggregated and variance is discussed openly.

5

Commercial Model

Weighted scores are cross-checked against real-world data (adoption, service metrics, cost per transaction).

6

Innovation and Risk

Stability, AI readiness, future-proofing

Another Approach



1

Content and NDC/Attribute Based Buying

Establish what matters most to your travelers and program — weighting inputs, not vendor defaults.

Why It Matters

Content varied **15-30%** between EDIFACT and NDC

Preferred OBTS **may not show their best negotiated fares** IF deals were NOT structured for the supported OBT channel

What We Measured

Identical itineraries tested across GDS EDIFACT, NDC Level 3, and direct connect — scored on fare availability, ancillary visibility, and seat map accuracy. NDC maturity level (1–4) and channel consistency across all booking entry points

What We Found

Suppliers with the best NDC marketing scored lowest on live content availability.

Demo environments are curated. Production is not. The gap was consistent across every supplier evaluated

“NDC-enabled” for domestic US ≠ NDC capability for international markets.

A critical gap for global programs that RFP responses never volunteered

Content scoring forced the conversations RFPs avoid:

Which airline / hotel programs are supported natively vs. via workarounds — and what written NDC / Attribute Based Buying roadmap commitments exist?

2

Servicing Impact

Assign weighted criteria across all six dimensions. Scores are 1–10; duty-of-care items carry a 2× multiplier

Why It Matters

Servicing cost — the fully-loaded expense of agent-assisted transactions, online booking failures, and exception handling — is the most underweighted dimension and has the **largest actual P&L impact**

A supplier winning on headline transaction fees can easily be the **most expensive option** once touchless booking rate, offline re-booking volume, and cost per agent interaction are factored in

What We Measured

Touchless booking rate against a defined itinerary complexity mix. Blended cost per assisted transaction including queue handling, re-ticketing, and reconciliation. Proactive (airline disruptions) rebooking SLA and after-hours fee structure

What We Found

Weighting servicing at **2× changed** the supplier ranking. Two programs on the same OBT showed a 34-point gap in touchless booking rate — driven by policy configuration, not the platform



UX and Adoption

Demos, RFP responses, reference calls, and live data capture — not marketing decks

Why It Matters

Consumer apps have permanently raised traveler expectations. Leakage results when travelers seek workarounds. Online adoption rate **directly drives managed travel cost savings** — yet it was rarely scored objectively. Demos are performance, not reality

What We Measured

Structured traveler feedback across mobile, desktop, and agent-assisted paths. Blind task completion testing across simple domestic, complex international, last-minute change, and policy exception scenarios. Policy transparency at point of purchase

What We Found

Identical demo performance, but 30 pt gap AT production

Assessment was carried out with genuine travelers operating in live environments governed by real policy rules

Lifting adoption by 10 points equates to annualized savings ranging from \$200K to \$500K

Purely from lower agent interaction costs, well before any fare savings enter the picture — it is fundamentally a P&L discussion.

Of every UX feature examined, presenting policy transparency at the point of booking proved the single highest-impact UX feature

Rather than hidden within a pop-up, it appears exactly when the decision is made, promoting compliance without reliance on training campaigns

4

Technology Stack

Evaluators score blind to prevent anchoring. Results are aggregated and variance is discussed openly

Why It Matters

Platform's technology stack impacts connection, automation, and reporting - now and in future. Integration architecture becomes a **strategic constraint**. The tech stack is key to creating value for all parties

What We Measured

API maturity: documented, versioned, production-grade access — not a “custom integration available upon request” footnote. Integration depth for expense, HR/profile, duty of care, and BI approvals. Data fidelity and latency: completeness of data elements and where gaps required manual reconciliation.

What We Found

“We integrate with XX OBT” means nothing without qualification. **API commercial terms** were a hidden differentiator — some charged per-call fees that materially affected total cost of ownership. Middleware dependency creates ambiguous support ownership during outages

5

Commercial Model

Weighted scores are cross-checked against real-world data (adoption, service metrics, cost per transaction)

Why It Matters

Differing pricing models (per trx, per-trip, subscription, hybrid, revenue-share) creates comparative challenges. Similar transaction fees can produce 40% different **total cost of ownership**. Incentive structures between buyers, TMCs, and OBT providers are often misaligned in ways not visible in the RFP

What We Measured

Normalized total cost of ownership: all fees translated into a common cost-per-trip equivalent using the program's actual transaction mix — applied consistently across all suppliers. Incentive transparency and embedded channel preferences. Contract risk: termination fees, volume commitment structures, and pricing adjustment mechanisms

What We Found

Opacity in pricing is itself a signal

Suppliers most willing to provide **fee transparency** upfront — without requiring NDA-level negotiation — were consistently the ones with the most defensible commercial models

Cost-per-trip normalization is non-negotiable.

It is the only valid comparison method when suppliers use different pricing structures. Without it you are comparing apples and oranges

Subscription models require scenario analysis, not point-in-time comparison

They looked expensive at low volume and dramatically cheaper at full program scale. The commercial evaluation **must model both**

6

Innovation and Risk

Stability, AI readiness, future-proofing

Why It Matters

A supplier who scores well today but is poorly positioned for the next three years is a liability, not an asset. M&A activity, new entrant disruption, and AI integration create market tension. Buyers who evaluated purely on current-state capabilities found themselves mid-contract with acquired platforms and broken integration roadmaps

What We Measured

Financial stability: revenue trajectory, investor backing, profitability trend, M&A signals. AI roadmap credibility: demonstrated **production** deployments with quantified outcomes. Innovation track record: what was deployed in 24 months vs. what was committed to in prior client conversations

What We Found

Every supplier led with AI in their RFP. **Fewer than half** could demonstrate a production feature set when asked directly. AI capability claims were the least correlated with actual delivery of any dimension evaluated. Reference checks on 24 months of roadmap delivery — not satisfaction scores — were the most predictive indicator of future risk

The New Evaluation Framework

Six dimensions that go beyond the traditional RFP

1

Content & NDC / Attribute Based Buying

True availability
across all
channels

2

Servicing Impact

Agent model,
touchless rate,
cost per tx

3

UX & Adoption

Booking
experience
scored by
traveler

4

Technology Stack

Integration
depth, API
maturity,
roadmap

5

Commercial Model

Pricing
transparency,
incentive
alignment

6

Innovation & Risk

Stability, AI
readiness,
future-proofing

Another Approach



What the Frameworks Revealed

+7%

Avg scoring gap between stated
Vs. measured NDC content parity

+30%

UX score variance across OBTs
on identical itinerary tasks

+12%

Higher CSAT compared to baseline

Key Learnings

- ✓ Content varies by live availability — demo environments $\neq$ production reality. Gap is availability, when applying price points, 15-30% gap.
- ✓ Weighting matters more than scoring. Two programs with the same scores reached opposite decisions because their priorities differed.
- ✓ Servicing cost is the most underweighted dimension in traditional RFPs and the one with the largest P&L impact.

Buyers: Take This Home

Your scoring framework — ready to apply immediately

1

Review the evaluation criteria you rely on today. Do they capture all six dimensions, UX, servicing impact, and innovation risk included?

2

Set your weightings before the RFP goes out. Let those weights express your program strategy rather than being worked backward from incumbent scores.

3

Score independently first, then talk through where the results diverge. Those disagreements reveal far more than the points on which everyone agrees.

4

Confirm your findings against actual data. Demos can mislead, whereas adoption rates, service costs, and traveler surveys will not.

Sellers: Take This Home

Your selling framework — ready to apply immediately

1

Refine your selling points so that all six dimensions come through, UX, servicing impact, and innovation risk among them

2

Play to your strengths while shoring up weaker areas, keeping both aligned with the client's program strategy

3

Showcase quick wins alongside longer-term gains, delivering the client tangible successes at every phase of the process

4

Test and retest to confirm that your demos hold up, thereby building trust and reducing the questions that might otherwise surface

Buyer and Seller Outcomes

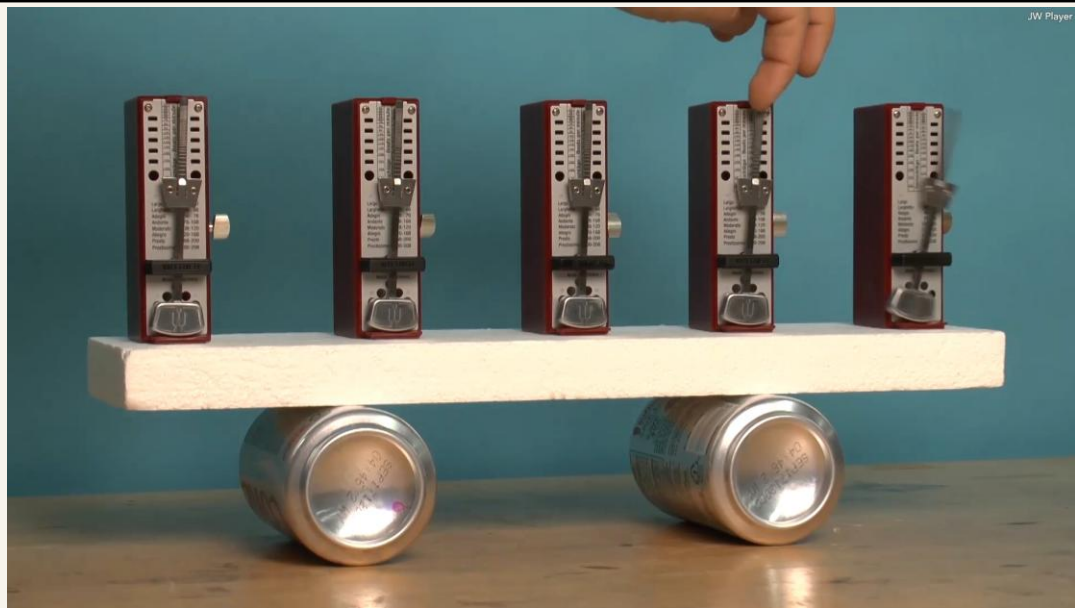
Less $\geq$ More

Strategy $\geq$ Accuracy

Direction $\geq$ Completeness

Traditional $\geq$ Nothing

Innovative $\geq$ Status Quo



About Us

Who We Are

Emily's Place provides a structured environment with direct support from a staff of house managers, case workers and counselors, and access to programs available in the community.

Emily's Story

When Mark Hagan settled on a name for the nonprofit he founded in 2002, he was honoring the memory of his cousin Emily Mays. In many ways, Emily's story echoes those of the women who find themselves at the place of hope named after her.

Our Team

Our team consists of leaders throughout our community who are dedicated to providing domestic abuse survivors with the practical training and education they need to establish stability in their lives and make positive choices going forward.

Transforming Lives. Building Communities. Restoring Hope.

[Donate Today](#)

Our Mission

emilysplacetx.org/donate/

Donate

The Evolving Travel Ecosystem

New Assessments for New Platforms

Will Tate | Partner, GoldSpring Consulting

will@goldspringconsulting.com