

THE TECH EDIT

HORWATH HTL · TECHNOLOGY PRACTICE

EDITION

03

JUNE 2026

DISCOVERABILITY & ASSET VALUE

Ranking, Referred, or Ignored? How AI Is Rewriting Hotel Discoverability

How generative search is moving demand from the search result to the model's answer, and what owners must do to keep an asset visible in the interfaces where guests now decide.



Ho Shyn Yee

Director & Technology Practice Lead

Ranking, Referred, or Ignored? How AI Is Rewriting Hotel Discoverability

By **Ho Shyn Yee**, Director & Technology Practice Lead

For twenty years, hotel discoverability has been treated as a marketing problem with a distribution answer. Rank on Google, buy the paid slots that matter, list on the OTAs, keep brand.com respectable, and the demand finds the hotel. Owners fund the machine, operators run it, and the mix of organic search, paid search, and OTA-referred traffic settles at some manageable share.

However, that machine has stopped working the way it used to. ChatGPT alone reached 800 million weekly active users by October 2025 on OpenAI's own disclosure, and every other major generative surface is scaling on a similar curve. The question a prospective guest asks — "best boutique resort in Bali with a beach club under USD 500" — increasingly gets answered inside the AI interface, and travellers' browsing and user journeys have changed accordingly.

The hotel's public interface used to be the search result. It is now the model's answer.

From discoverability to answerability

Three generations of hotel discoverability sit inside the last twenty years, and each has changed the job the property has to do.

The **SEO era**, running from the early 2000s through the middle of the last decade, made discoverability the discipline. Hotels invested in keyword-rich content, backlink profiles, and site architecture in order to earn organic rank on search engines such as Google. The output was traffic, and traffic was the currency of demand.

The **SEM era** overlapped and then dominated. Google's economics tilted from discoverability to paid visibility; the top of the results page filled with sponsored slots that owners, operators, and OTAs all bid on. The game shifted from earning traffic to buying traffic. OTAs entered the auction with deeper pockets and smarter algorithms, direct channels ceded ground, and the industry settled into a distribution mix in which a large share of demand arrived through intermediaries.

The **GEO (or AEO) era**, where the industry now sits, is different. When a traveller asks a generative engine — ChatGPT, Claude, Perplexity, Gemini, Google's AI Overviews, Copilot — for a hotel recommendation, the interface returns a synthesised answer with named properties and personalised reasoning, not a list of blue

links. On its Q3 2025 earnings call, Tripadvisor's chief executive attributed ongoing declines in referral traffic to the changing search landscape and the rise of AI overviews — one of the first listed-company acknowledgements that AI answers are materially eroding travel referral traffic. Deloitte's 2025 US Holiday Travel Survey traced the underlying shift: generative-AI use in trip planning moved from 8% in 2023 to 16% in 2024 to 24% in 2025.

Asia-Pacific tends to lag on some technology-adoption curves and lead on others. Broader consumer AI interest across Southeast Asia runs high on the available public data, and regional travellers are unlikely to sit out the shift. Hotel-specific AI-adoption data at the Asia-Pacific level, however, remains limited; the observations that follow apply directionally across the region rather than being drawn from any single market study.

Three factors that determine AI legibility

Three key factors influence whether a large language model cites a property or ignores it — and all three are within the property's control.

CONTENT DEPTH

The model has to cite something, and it prefers to cite the source that answers the question. On the hotel's own domain, that means substantive guides, editorial depth on the destination, and structured answers to the specific questions guests are likely to ask. Where a property's website is thin, the model reaches instead for the OTA or the third-party editorial piece, and the hotel loses control of how it is described. Across the properties we have audited in Asia-Pacific, OTA and third-party editorial content routinely dominate model citations, while brand.com trails by a wide margin. The gap is a content-depth problem more than a distribution one.

SIGNAL QUALITY

Models typically verify claims across multiple sources before including a property in an answer. That verification depends on the consistency of how the property is described across the web — name, address, category positioning, brand descriptors — and on the presence of independent third-party authority, from editorial mentions to review coverage. Where signals conflict, the model routes around the property. Where signals align, the model has a basis for citation.

TECHNICAL ACCESSIBILITY

The last factor is the one most often overlooked. If the model cannot access the site, nothing else matters. Schema markup (JSON-LD, hotel schema) tells the model what type of asset it is looking at. Robots.txt settings determine whether AI-specific crawlers — including GPTBot, ClaudeBot, Google-Extended, and PerplexityBot — can read the property's content at all. AI-crawler traffic has grown by orders of magnitude in the last two years across the sites we monitor, and access is now the gating condition. A property that has locked AI crawlers out of its site cannot be recommended by the models.

Where the transition breaks

Four mistakes recur in the properties and portfolios navigating this shift in user behaviour, and each has a booking consequence that is critical to understand.

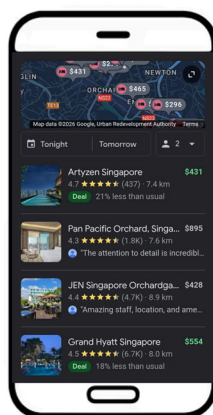
Treating GEO as SEO with more keywords. The keyword-density instinct does not apply to GEO. Models read for meaning, not for term matching, and content built around keyword frequency reads to the model as low quality. In our audit work across Asia-Pacific properties, we routinely see pages that rank well on Google fail to appear in AI-answer citations for the same query, and pages absent from the top of Google surface in the AI answer. Ranking well on Google and being cited in the AI answer are increasingly different disciplines.

Betting on AI-generated content for scale. Models increasingly detect and deprioritise their own kind. Hotel websites filled with AI-written editorial in the name of content velocity are finding that the volume now works against citation rather than for it. Original, substantive content still wins, and it is expensive to produce well.

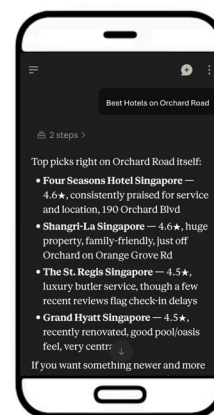
Assuming OTA presence covers AI legibility. An OTA listing tells the model the property exists; it does not, in most cases, give the model the primary-source depth it needs to include the property in a specific-answer query. Being on Booking.com is not the same as being in the answer to "best boutique hotels in Ubud with pool villas under USD 400."

Neglecting the direct website as the primary source of truth. The model has to cite something. When brand.com or the property website is thin, generic, or paywalled, the model reaches for whatever else exists — an OTA, an editorial roundup, a review site — and the hotel's description of itself is replaced by someone else's.

GOOGLE SEARCH



CLAUDE



Same search, two interfaces: Google returns a map and price-ranked OTA listings for "best hotels on Orchard Road," while Claude returns a synthesised answer naming hotels directly.

What this means for capital decisions

For owners, asset managers, and senior operating teams reviewing distribution, brand, and technology investments over the next 12 to 24 months, four questions are worth putting to the operating team.

1. How discoverable is the asset in AI-answer surfaces today?

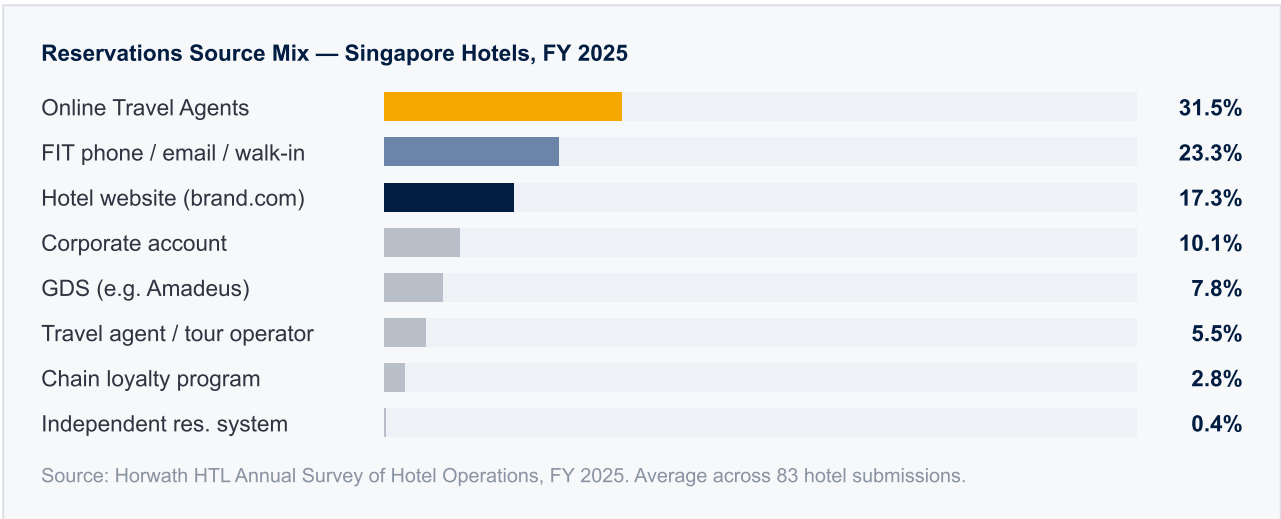
The exercise is straightforward. Assemble a realistic set of guest queries covering comparison, category, and destination questions, and run them across ChatGPT, Claude, Perplexity, Gemini, Copilot, and Google's AI Overviews. Catalogue whether and how the property surfaces. Most owners and operators have never done this, and are surprised by the answer. In our own diagnostic scans across Asia-Pacific markets, most of the properties we audit do not surface in AI-answer sets for their core comparison queries, and the ones that do are concentrated among a small number of large brand groups. A property missing from the answer is a property not competing in that channel.

2. Who controls the content and channels that determine that discoverability?

For branded properties operating under a management agreement, the question surfaces the same contractual gap this series has previously flagged: who owns the brand.com content, the property's editorial voice, the schema markup, and the AI-crawler permissions? If the answer is "the brand," the owner's ability to influence AI-answer visibility is limited to what the brand chooses to prioritise. For independent properties, the question is where the primary-source content lives, and whether it is thin or substantive.

3. What is the exposure if AI-mediated discovery displaces OTA-mediated discovery?

D-Edge's 2024 distribution analysis put the Asia-Pacific hotel channel mix at roughly 68% OTA and 32% direct across the region, with APAC luxury properties reaching 56% direct. Skift Research's late-2024 outlook forecast direct bookings overtaking OTA bookings globally by 2030. However, based on our proprietary Annual Survey of Hotel Operations FY 2025 in Singapore, for example, OTAs account for 32% of all reservations, followed by 23% by phone or email, and 17% by brand.com.



If a meaningful share of top-of-funnel discovery moves from Google to AI, and if AI-answer citations continue to favour brand.com content over OTA listings, a hotel over-concentrated on OTA distribution is exposed twice: once at the front of the funnel, and once in commission economics. Voice-based AI agents — advanced LLM phone assistants — can be used to manage phone and email bookings, keeping the framework fully exhaustive.

4. Is our current digital investment supporting AI legibility, or working against it?

The audit here is unglamorous but necessary. Thin marketing copy, keyword-density SEO, AI-authored blog content, and paywalled experience information all work against citation. Substantive editorial, clean schema, permissive crawler policies, and consistent third-party signal all work for it. Where we see AI-referred visitors arriving at Asia-Pacific hotel sites today, the volume is small but their intent is measurably higher — deeper page engagement, longer sessions, and lower bounce than search-referred visitors. The economic weight of that traffic will lag its behavioural weight, but only for so long.

Risks worth considering

Three risks are worth considering before an owner reallocates spend against this shift.

MISREPRESENTATION RISK

Models can and do misrepresent properties in their answers — pulling outdated descriptions, misstating category, or conflating a property with a similarly named one nearby. AI legibility does not, on its own, guarantee accurate representation. Owners investing here should invest in equal measure in monitoring what the models are saying, not only in whether the property surfaces at all.

MEASUREMENT IMMATURITY

The tools for measuring AI-answer citation are new, inconsistent, and largely built by vendors with their own commercial interests. Two audits of the same property can produce materially different visibility scores depending on the query set, the models tested, and the sampling window. Any single measurement is a directional signal, not a definitive score.

COMPETITIVE ARMS-RACE RISK

As more properties invest in GEO/AEO, the cost of maintaining a top-tier position rises — mirroring the arc of paid search. Owners moving early may capture disproportionate return; those catching up two years from now may find that the incremental cost of visibility has expanded to look like the paid-search economics they were trying to escape.

A pragmatic response

Owners and operators moving on this now have a modest set of actions to prioritise. Here are the top three worth putting on the agenda.

Firstly, run an AI-answer visibility audit as a quarterly baseline rather than an annual review. Models change, citation profiles change with them, and a property's position in the answer set drifts over quarters, not years. The first audit establishes where the asset stands today; the discipline of running it quarterly is what surfaces drift before it becomes a booking gap. With AI tools, this audit can be automated with pre-built skills, and the appropriate reporting, escalation, and actions established.

Secondly, rebuild direct-website content around substantive answers, with technical accessibility treated as ongoing hygiene rather than project work. Guides that answer real guest questions, and editorial depth on the destination. Schema markup, crawler permissions, and consistent third-party signal should be maintained as a discipline. This discipline used to feed Google rank; today its return shows up primarily in whether the model can cite the property when the guest asks.

Thirdly, put the question to the operator directly, since the search agency is unlikely to have the authority to change most of it: what is the specific plan for content depth, schema markup, and AI-crawler permissions, who owns it, and on what cadence is it reviewed. Many management teams can produce an SEO report on request; few can say with confidence whether GPTBot, ClaudeBot, Google-Extended, and PerplexityBot are even permitted onto the property's site.

An asset-visibility audit gives an owner the data needed during an HMA renegotiation or brand conversion to challenge a brand's marketing fee justification if the brand is failing to maintain AI legibility.

How we can help

The Horwath HTL Technology Practice advises hotel owners, asset managers, and investors on the technology decisions that shape asset value. We are independent of vendors, and our recommendations are accountable only to the client's investment case. Our **Diagnostic Engagements** provide the AI-answer visibility baseline this article argues for — a structured audit of where an asset or a portfolio currently stands across the major generative and answer engines, and which of the three underlying factors is the binding constraint. Our **Tech Edit LAB In-House** programme takes senior leadership teams through the shift from search to answer at the level of capital and distribution decisions, where the choices actually get made.

If you are reviewing your distribution mix, planning an HMA renewal or brand conversion, or looking at what your asset's discoverability profile looks like, we would be glad to have a conversation.



Ho Shyn Yee · Director & Technology Practice Lead

syho@horwathhtl.com · The global leader in hospitality, tourism and leisure consulting