

The Grounded Front Desk

Guest communication, response latency and factual reliability — what the evidence supports, what it does not, and what a hotel should actually require of an AI system.

RESPONSE LATENCY, TO SCALE

7.9 seconds

measured average first response, two-hotel pilot

average response time among companies that responded at all, across 2,241 firms audited (HBR, 2011)

42 hours

The bar above is not to scale. At true scale it would run roughly 19,000 times the width of the line above it.

FOUR NUMBERS THAT FRAME THE DECISION

4.2

full-time staff required for
single-person 24/7 cover

23%

of audited companies never
answered the enquiry at all

22–94%

hallucination rate range across
26 current models, AI Index 2026

424k

guest messages handled across
two hotels in ~7 months

ONE CORRECTION THIS REPORT MAKES TO ITS OWN INDUSTRY

What the industry says

"Replying within an hour makes you seven times more likely to convert." Repeated in almost every hotel-technology article we checked.

What the source actually says

Seven times more likely to *qualify* a lead than firms replying **one hour later** — not than everyone else. Qualification, not conversion. §4 explains.

The real finding is still strong enough to build on. It just is not the one being quoted.

Research report • September 2026 • Data current to Q3 2026

Scope: Guest enquiry handling in hotels — coverage economics, response latency evidence, factual reliability of language models, legal exposure, and a deployment specification

Sources: 53 cited references including peer-reviewed work from *Marketing Science*, *Journal of Consumer Research*, *Manufacturing & Service Operations Management*, *ACM Computing Surveys*, *Tourism and Hospitality*, and the proceedings of NeurIPS, ICLR, ACL and EMNLP — alongside Eurostat, the US Bureau of Labor Statistics, the Turkish Ministry of Labour and Social Security, Stanford HAI, NIST, EUR-Lex, HOTREC, the American Hotel & Lodging Association, and a decision of the British Columbia Civil Resolution Tribunal

Published by: Vera — see the sourcing, dataset and independence statement in §2

CONTENTS

What this report covers

- 1** **Executive summary**
Eight findings, one industry claim corrected, and one thing we could not measure
- 2** **Method, sourcing and the pilot dataset**
What we measured, how, over what period — and what the dataset does not establish
- 3** **The coverage problem, not the headcount problem**
Why 24/7 availability is arithmetic and throughput is an argument
- 4** **What the response-time evidence actually says**
The real Harvard Business Review finding, and the gap nobody has filled
- 5** **The labour picture, honestly**
Including the part where the shortage has eased
- 6** **Why generic assistants fail at the one question that matters**
Parametric generation versus a live system of record
- 7** **Hallucination is not a bug you can patch out**
What the current measurements and the formal results show
- 8** **The liability that arrived in 2024**
Moffatt v. Air Canada, and the duty it created
- 9** **The pilot: two hotels, seven months**
The dataset, its definitions, and its limits
- 10** **The economics, on the conservative basis**
Costed at the top tier and heaviest usage, not the cheapest plan
- 11** **A deployment specification**
Ten requirements, each traced to a cited finding
- 12** **Limitations and data gaps**
Eleven things this report cannot establish
- 13** **References & resources**

WHY A VENDOR IS PUBLISHING THE COUNTER-EVIDENCE

This report is published by a company that sells AI guest communication software. It contains a section on peer-reviewed research finding that booking intention is *lower* with a chatbot than with a human (§7), a randomised field experiment showing that disclosing a bot's identity cut purchase rates by close to eighty per cent (§6), and a limitations section listing what our own pilot data cannot prove (§12).

That is the point. Every operator reading this has already been sent a dozen decks of favourable statistics. The useful document is the one that tells you where the evidence runs out — because that is where your risk is.

SECTION 01

Executive summary

A hotel that wants to answer guest enquiries at three in the morning has an arithmetic problem before it has a technology problem. Everything in this report follows from that arithmetic, and from one legal decision handed down in February 2024.

4.2

Full-time staff required for single-person 24/7 cover, before absence or peak load

€15,215

Monthly labour cost of that cover at EU average hospitality labour cost^{5,6}

23%

Of 2,241 audited companies that never answered a web enquiry at all¹

22–94%

Hallucination rate range across 26 current models on AA-Omniscience³²

Eight findings

1 — Round-the-clock coverage costs 4.2 full-time staff before anyone asks how fast they work. A week contains 168 hours; a full-time post covers 40. That ratio is arithmetic, not a productivity estimate, and it holds whether the desk receives ten enquiries a day or ten thousand. **AUTHOR'S CALCULATION**

2 — At official European labour statistics, that coverage costs between €9,391 and €15,215 a month. Eurostat puts 2025 hourly labour cost in accommodation and food service activities at €20.9 across the EU-27, €19.0 in Greece, €17.6 in Spain and €12.9 in Portugal.⁶ Applied to the 728 coverage hours in an average month, those are the monthly figures. **OFFICIAL**

3 — The response-time statistic the industry quotes is not what the source says. The Harvard Business Review study behind it found firms contacting a lead within an hour were nearly seven times as likely to *qualify* that lead as firms contacting it **one hour later** — and more than sixty times as likely as firms waiting twenty-four hours or more.¹ The comparison is one hour against two, the outcome is qualification rather than sale, and the data is 2011 cross-industry sales, not hospitality. Section 4 sets out what can honestly be built on it.

4 — No peer-reviewed research links hotel enquiry response time to booking conversion. We searched for it specifically. Every result was a hotel-technology vendor recycling the same 2011 figures. This is a real hole in the literature, and the transfer of cross-industry sales findings to hotel enquiries is an assumption rather than a demonstrated result. We say so rather than hide it.

5 — The labour shortage has eased in the United States and persists in Europe. US hospitality job openings fell from an annual average of 9.5 per cent in 2022 to 5.6 per cent in 2025, and by July 2026 the sector's openings rate (4.5 per cent) was barely above the all-industry rate (4.4 per cent).^{11,12} What remains durable is turnover — the quits rate in accommodation and food services runs about 1.8 times the all-industry rate — and the European vacancy gap, where hospitality sits 0.9 percentage points above the whole-economy average.^{10,11} **OFFICIAL**

6 — Hallucination is not a defect that better models remove. Stanford's AI Index 2026 reports hallucination rates ranging from **22 to 94 per cent** across 26 current models on the AA-Omniscience benchmark.³² A 2024 result in the *Proceedings of the ACM Symposium on Theory of Computing* establishes "an inherent statistical lower-bound on the rate that pretrained language models hallucinate certain types of facts, having nothing to do with the transformer LM architecture or data quality."²⁷ **PEER-REVIEWED** That bound applies to facts appearing once in training data. Tonight's rate for a specific room type is exactly such a fact.

7 — Since February 2024 there has been a legal standard, and it is an ordinary negligence standard. In *Moffatt v. Air Canada* the British Columbia Civil Resolution Tribunal rejected the airline's argument that it could not be liable for its chatbot, calling the submission that the chatbot was a separate legal entity "a remarkable submission," and found that "Air Canada did not take reasonable care to ensure its chatbot was accurate."⁴¹ The award was trivial; the doctrine is not.

8 — Disclosure becomes mandatory in the EU on 2 August 2026, and disclosure has historically cost conversions. Article 50 of the EU AI Act requires that people be informed they are interacting with an AI system.⁴³ A randomised field experiment across more than 6,200 customers, published in *Marketing Science*, found that disclosing chatbot identity before the conversation "reduces purchase rates by more than 79.7 per cent" — and that the mechanism was customers perceiving the disclosed bot as *less knowledgeable* and *less empathetic*.⁴⁴ **PEER-REVIEWED** Once disclosure is compulsory, being demonstrably knowledgeable is the only lever left.

What the pilot found, and what it does not prove

Across two hotel properties under common ownership, running continuously for approximately seven months, the system described in this report handled roughly **424,000 guest messages**. In a representative 30-day window it handled **60,611** — an average of 2,020 a day, or 84 an hour around the clock — resolving **99 per cent** without human intervention, ending **64.8 per cent** of conversations with a direct booking link, at an average first response of **7.9 seconds**. Every price quoted was checked against the hotels' live reservation system before being sent.

What this does not establish. There is no control group, no revenue attribution, no guest-satisfaction instrument, and two properties under one ownership is not a

sample. It is an existence proof that the architecture works at volume, not evidence of what it is worth. Section 9 gives the full definitions and Section 12 the full limitations.

Method, sourcing and the pilot dataset

Sourcing rules

One — official statistics for anything about money or labour.

Wage and labour-cost figures come from Eurostat, the Turkish Ministry of Labour and Social Security, and the US Bureau of Labor Statistics. No trade-press summary is used where the underlying series is published.

Two — the original paper, never the summary of it. Behavioural and technical findings are cited to the journal or conference proceedings that published them, with sample sizes stated where the source discloses them, and marked **VENDOR-REPORTED** where a commercially interested party published the number.

Three — the counter-evidence appears at full strength.

Research finding that chatbots reduce booking intention, that disclosure suppresses purchases, and that current agents fail more than half of realistic customer-service tasks is in the body of this report, not a footnote.

Four — costs are modelled on the most expensive assumption, not the cheapest. The economics in §10 use the publisher's top-tier plan and an AI usage estimate computed for intensive use: long, detailed prompts, a capable model, and heavy tool calling. Most vendor ROI arithmetic runs the other way. Ours is deliberately unflattering so that a reader's real cost is more likely to fall below it than above.

The pilot dataset

The operational figures in this report are the publisher's own first-party data. They are presented with their measurement frame stated so that a reader can judge what they support.

EVIDENCE CLASSIFICATION

PEER-REVIEWED

Journal or refereed conference proceedings.

OFFICIAL

Statistics agency, ministry, regulator or tribunal.

INSTITUTIONAL

Standards body, trade association or university institute.

CALCULATION

Derived by the authors from cited inputs, shown in full.

VENDOR-REPORTED

Published by a commercially interested party.

FIRST-PARTY

The publisher's own operational data, with its measurement frame stated.

ON CURRENCY

Amounts are quoted in the currency of the source. Labour costs are in euro because Eurostat publishes them that way; Turkish figures are in lira from the Ministry; platform costs are in US dollars from the published price list. No exchange-rate conversion is performed anywhere in this report, because a single converted figure would imply a precision that currencies moving at this speed do not support.

A NOTE ON THE SEVEN-MONTH FIGURE

The 424,000 cumulative total is the measured 30-day volume extrapolated across the deployment period. It assumes a stable monthly run rate, which seasonality in a resort market makes unlikely to hold exactly. It is given as an order of magnitude and rounded accordingly.

Parameter	Value
Properties	Two hotels under common ownership, not named at the operator's request
Time in service	Approximately seven months of continuous live operation
Detailed measurement window	One 30-day period within that span
Messages in window	60,611 guest messages
Cumulative volume	Approximately 424,000 messages — the 30-day figure extrapolated across seven months, and stated as approximate for that reason CALCULATION
"Message"	One inbound or outbound message in a guest conversation, across web chat, WhatsApp and Instagram
"Conversation"	One continuous guest session; the booking-link rate is expressed against conversations, not messages
Autonomy (99%)	Share of enquiries closed without a human agent taking over the conversation
Booking-link rate (64.8%)	Share of conversations ending with a generated, guest-specific booking link
Response time (7.9s)	Average time to first substantive reply, measured in the window
Grounding	Rates and availability read from the hotels' live reservation system through an authenticated integration layer before any price is sent

WHAT THE DATASET DOES NOT ESTABLISH

No control group, so no causal claim about bookings. No revenue attribution, so no return-on-investment figure. No guest-satisfaction instrument, so nothing about experience quality. Two properties under one ownership, so no generalisation to other operators, segments or markets. No independent audit. Read it as an existence proof at volume, and nothing more.

Independence statement

This report is published by Vera, which sells the category of software it examines. No competitor is cited as a source of fact. The sections most likely to inform a purchasing decision — §7, §12 — carry the evidence least favourable to the publisher, and the pricing used in §10 is the published list price, which any reader can check.

The coverage problem, not the headcount problem

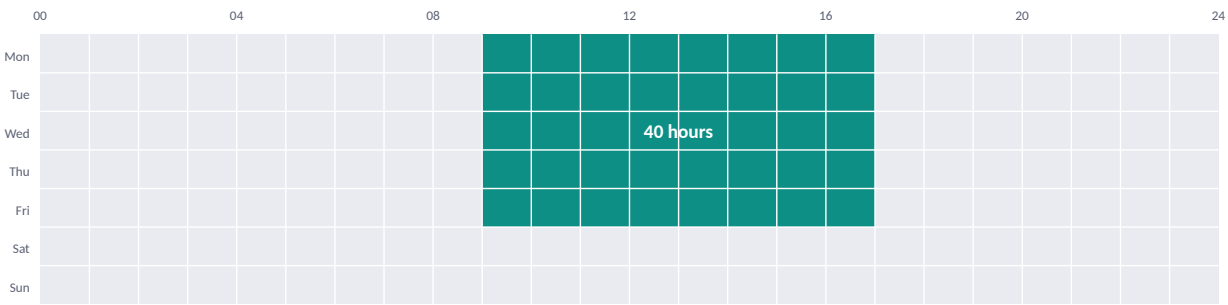
Most arguments for automating guest communication start with how many messages a person can handle. That is the weakest available argument, because nobody has credibly measured it. The stronger one needs no productivity assumption at all.

A week contains 168 hours. A full-time post covers roughly 40 of them. To have **one** person available at every hour of every day therefore requires **4.2 full-time posts** — before annual leave, before sickness,

before breaks, and before any allowance for two enquiries arriving at once. **AUTHOR'S CALCULATION** This holds whether the property receives ten enquiries a day or ten thousand.

One week, 168 hours — and what one full-time post covers

Each cell is one hour. Shaded cells are a single Monday-to-Friday, 09:00–17:00 post.



168 ÷ 40 = 4.2

full-time posts for one person on duty at all times

Before any of the following:

- annual leave • sickness • breaks • training • two enquiries at once
- a second language • a third language

Figure 1 — The coverage arithmetic. **AUTHOR'S CALCULATION** from a 168-hour week and a 40-hour full-time post. No productivity, throughput or message-volume assumption is used. The shaded block is illustrative of one conventional shift pattern; the ratio is unaffected by how the hours are arranged.

What that coverage costs

An average month contains **728** coverage hours (168 × 52 ÷ 12). Eurostat publishes hourly labour costs by economic sector; for accommodation and food service activities in 2025 the figures are €20.9 across the EU-27, €19.0 in Greece, €17.6 in Spain and €12.9

in Portugal.⁶ **OFFICIAL** Labour cost, in Eurostat's definition, is compensation of employees plus taxes minus subsidies — that is, what the employer actually pays, not what the employee receives.

Monthly labour cost of one person on duty, around the clock

728 coverage hours × Eurostat 2025 hourly labour cost, accommodation and food service activities.

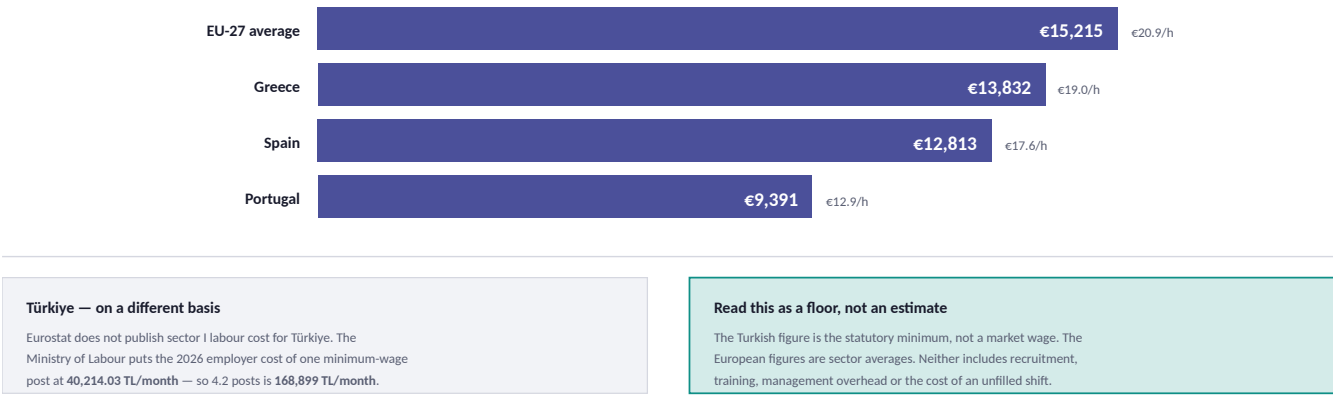


Figure 2 — Cost of round-the-clock coverage. Hourly labour costs: Eurostat dataset `lc_lci_lev`, NACE Rev. 2 Section I, 2025.⁶ Turkish employer cost: T.C. Çalışma ve Sosyal Güvenlik Bakanlığı, 2026 minimum wage schedule.⁷ Monthly totals are **AUTHOR'S CALCULATION** at 728 coverage hours and 4.2 posts respectively. Currencies are not converted.

THE THROUGHPUT CLAIM WE ARE DELIBERATELY NOT MAKING

Almost every document in this category converts message volume into a headcount — "1,700 enquiries a day requires 7 to 9 staff" — on the basis of an assumed messages-per-hour rate. We could not find a single credible independent source for such a rate. The one figure in general circulation, 274 chats per agent per month, is attributed by the trade publication carrying it to a chat-software vendor, and implies fewer than two conversations an hour.¹⁷ **VENDOR-DERIVED**

The peer-reviewed literature on contact-centre operations is clear about the cost structure — agent salaries "typically account for 60–70% of the total operating costs" — and about what large, well-run operations achieve, with "agent utilization levels ... between 90% to 95%."¹⁶ **PEER-REVIEWED** But those utilisation economics come from *scale*, through queueing effects a single hotel front desk cannot access. So this report makes no headcount-equivalence claim anywhere. The coverage arithmetic above needs none.

What the response-time evidence actually says

One 2011 study underwrites almost every claim made about speed of reply in this industry. It is a good study. It does not say what it is quoted as saying.

Oldroyd, McElheran and Elkington, writing in *Harvard Business Review*, audited how 2,241 companies responded to test web enquiries and separately analysed 1.25 million sales leads across 42 firms.^{1,2} The finding, verbatim: firms that tried to contact potential customers within an hour of receiving a query "were nearly seven times as likely to qualify the lead ... as those that tried to contact the customer **even an hour later** — and more than

60 times as likely as companies that waited 24 hours or longer."

Three things follow that the common paraphrase loses. The seven-fold comparison is **one hour against two hours**, not one hour against everyone else. The outcome measured is **qualification** — defined in the paper as having a meaningful conversation with a key decision maker — not a sale. And the sixty-fold figure, which is the genuinely dramatic one, is the comparison against a day's delay.

How 2,241 companies actually responded to a web enquiry

Harvard Business Review audit, 2011. Percentages are of all audited companies.



WHAT THE ODDS COMPARISON IS — AND IS NOT

~7x more likely to **QUALIFY** a lead

Contacting within 1 hour, compared with contacting **one hour later** — a two-hour reply, not a slow one.

"Qualify" = a meaningful conversation with a decision maker. Not a booking.

>60x more likely to **qualify** a lead

Contacting within 1 hour, compared with waiting **24 hours or more**. This is the comparison that matters for an enquiry arriving overnight — and it is the one least often quoted.

The average among companies that responded at all: **42 hours**

Not the slow tail. The average. Nearly a quarter of audited companies never replied to the enquiry at any point within 30 days.

The opportunity in this data is not being fast. It is being present at all.

Figure 3 — The response-time evidence, stated precisely. Source: Oldroyd, McElheran & Elkington, "The Short Life of Online Sales Leads", *Harvard Business Review* 89(3), March 2011.^{1,2} Audit sample 2,241 US companies; separate lead analysis covered 1.25 million leads across 29 B2C and 13 B2B firms. The study is cross-industry sales, not hospitality — see the caution below.

The related study that is not what it is called

A second body of figures circulates as "the MIT study" — most often the claim that the odds of contacting a lead fall a hundredfold between five minutes and thirty minutes.³ That research does exist, it covers more than 15,000 leads and 100,000 call attempts across six companies, and its findings are as stated. But it is **not an MIT publication**. It is platform data from a sales-software company, analysed by a researcher who held an MIT affiliation at the time. **VENDOR-REPORTED** A later audit by the

same company, testing 9,538 businesses, found 47 per cent gave no response at all and the median first phone response ran to three hours and eight minutes.⁴

THE GAP THIS REPORT WILL NOT PAPER OVER

We searched specifically for peer-reviewed research linking *hotel* enquiry response time to *booking* conversion. We did not find any. Every result was a hospitality-technology vendor recycling the 2011 figures above. The speed-to-lead evidence base is cross-industry B2B and B2C sales, and its transfer to a guest asking about a room rate on WhatsApp is an assumption, not a demonstrated result. It is a reasonable assumption — the mechanism is the same — but a reader is entitled to know it has not been tested in this setting, and any document that presents it as hospitality research is overstating the literature.

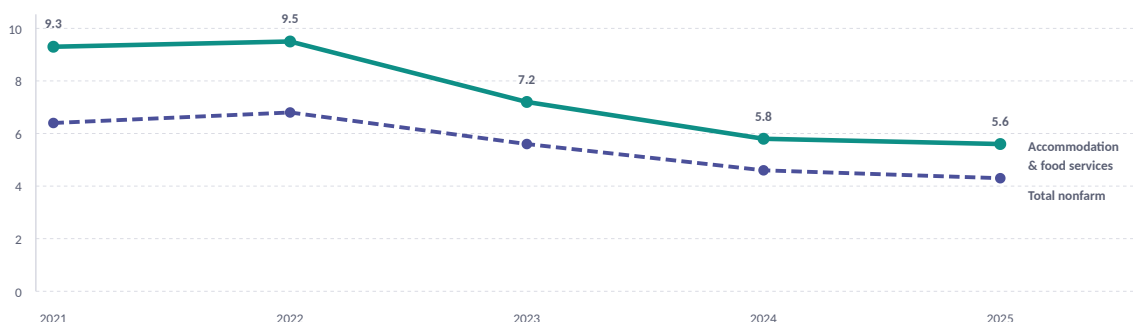
SECTION 05

The labour picture, honestly

The hospitality staffing crisis is real in Europe and substantially eased in the United States. A report claiming an ongoing acute shortage across both would be contradicted by current official data, so this one does not.

US job openings rate: hospitality has converged on the whole economy

Annual averages, per cent. Accommodation and food services against total nonfarm.



WHAT HAS NOT CONVERGED

US quits rate, July 2026

Accommodation & food **3.5%** vs all industry 1.9% — 1.8x

EU job vacancy rate, Q1 2026

Hospitality **3.2%** vs whole economy 2.3% — +0.9pp

Figure 4 — The labour picture. US annual average job openings rates: Bureau of Labor Statistics, JOLTS Table 16.¹² Quits rates: JOLTS Table 4, July 2026 preliminary.¹¹ EU vacancy rates: Eurostat, Q1 2026, NACE Section I.¹⁰ **OFFICIAL**

The European position is different. HOTREC's January 2026 assessment puts the EU hospitality sector at **10 per cent short of its workforce** across a sector employing ten million people in two million

businesses, ninety per cent of them micro-enterprises with fewer than ten staff.¹⁵ **INSTITUTIONAL**

Those figures come from national member associations rather than a harmonised statistical

survey, so they should be read directionally — but the Eurostat vacancy gap corroborates the direction.

In the United States, the American Hotel & Lodging Association's February 2026 survey of 246 hoteliers found more than half reporting their properties somewhat or severely understaffed, with labour costs named a top financial challenge by 65 per cent.¹³ An earlier wave gives the only role-level figure we located: asked which positions they could not fill, operators named housekeeping first at 38 per cent and **front desk second at 26 per cent.**¹⁴

WHY TURNOVER MATTERS MORE THAN VACANCIES FOR THIS DECISION

A vacancy is a cost you can see. Turnover at roughly twice the all-industry rate is a cost you cannot: every departure takes with it the knowledge of which room types connect, which rate plan applies to a child of eleven, and what the cancellation policy actually says in March. An automated layer does not resolve turnover. What it does is hold that knowledge in a system rather than in a person — provided the system reads it from the reservation database rather than from its own memory, which is the subject of the next three sections.

SECTION 05, CONTINUED

Where guests now ask

The coverage problem in §3 would be manageable if enquiries still arrived by telephone during office hours. They do not. DataReportal's global tracking puts WhatsApp at **54.4 per cent of online adults** using it monthly, against 5.79 billion social media user identities worldwide.¹⁹ **INSTITUTIONAL** In Türkiye specifically, 88.3 per cent of the population are internet users, 70.9 per cent hold social media identities and mobile connections run at 93.3 per cent of population.²⁰

The volume of *automated* business messaging is moving faster than any survey can track. On its Q1 2026 earnings call Meta disclosed "more than 10 million conversations each week being facilitated through Business AIs, up from 1 million at the start of the year" — a tenfold increase in a single quarter.²¹

COMPANY DISCLOSURE Whatever a hotel decides, its guests are being trained on this interaction pattern elsewhere.

54.4%

Of online adults using WhatsApp monthly, globally¹⁹

1m → 10m

Weekly Meta Business AI conversations, within one quarter of 2026²¹

98% / 32%

Hotel owners who have begun using AI, versus those with it embedded²⁴

Guest appetite has been surveyed, though the best available figure comes from a software vendor and should be read as such: research published by Oracle Hospitality with Skift, covering 5,266 consumers and 633 hotel executives across nine markets, reported that **77 per cent of travellers are interested in using automated messaging or chatbots for customer service requests at hotels.**²² **VENDOR-REPORTED**

Operator adoption is broad and shallow. A 2026 franchisor survey found 98 per cent of hotel owners had begun incorporating AI while only 32 per cent said it was embedded across most of their

operations, and 73 per cent wanted to do more but felt overwhelmed.²⁴ **VENDOR-REPORTED** The sample size is not disclosed. Separately, the *Lodging Technology Study* put hotel IT budgets at more than 4.2 per cent of revenue, with 76 per cent of respondents naming increasing employee productivity among their top technology priorities.²³ **INSTITUTIONAL**

And the guests themselves arrive speaking other languages. Türkiye received 52.8 million foreign visitors in 2025, with Russia, Germany and the United Kingdom the three largest source markets at 6.90, 6.75 and 4.27 million respectively — three

languages, none of them Turkish.⁵⁰ Spain, for comparison, received 96.8 million international tourists in 2025, led by the United Kingdom, France

and Germany.⁵¹ **OFFICIAL** A single overnight shift covering Russian, German and English is not a staffing plan; it is three staffing plans.

SECTION 06

Why generic assistants fail at the one question that matters

A hotel guest asks one question more than any other, and it is the single question a language model is least equipped to answer from its own memory: what does it cost.

A price is not general knowledge. It is a fact that exists in one place, changes daily, depends on dates, occupancy, child ages, minimum-stay rules and which rate plan applies, and has no correct answer outside the reservation system that holds it. A model asked for it without access to that system will still produce a number, because producing plausible text is what it does.

The distinction has a name in the research literature. Lewis and colleagues, introducing retrieval-augmented generation at NeurIPS 2020, framed it

precisely: language models' "ability to access and precisely manipulate knowledge is still limited," and "providing provenance for their decisions and updating their world knowledge remain open research problems."³³ **PEER-REVIEWED** Rates change nightly. A model relying on stored parameters is stale by construction, and cannot say where its answer came from. A subsequent survey of the field frames the same three failure modes as "hallucination, outdated knowledge, and non-transparent, untraceable reasoning processes."³⁵

The same question, two architectures

"Family room, three nights in July, two adults and a child aged nine — how much?"

UNGROUNDED — ANSWERS FROM PARAMETERS	GROUNDED — ANSWERS FROM THE RECORD
1 • Model reads the question Intent and entities extracted correctly. This part works.	1 • Model reads the question Identical step. The difference begins at step two.
2 • Generates a number from training data No rate table exists in the model. It produces what a rate for a hotel like this would <i>plausibly</i> look like.	2 • Calls the reservation system An authenticated query returns live availability, the applicable rate plan, minimum stay and the child-age policy.
3 • No provenance, no check Nothing to compare the answer against, so nothing can flag it as wrong — not the system, not the guest.	3 • Answer constructed from what was returned If the record does not cover the case, the correct output is a clarifying question — not a number.
4 • Guest receives a price the hotel never set Minimum-stay rules, child-age brackets and multi-room splits are invented along with it.	4 • Guest receives the hotel's own price Plus a booking link encoding the same parameters, so the quote and the checkout cannot diverge.

The models are the same. The architecture is the whole difference — and it is the part a buyer can actually specify.

Toolformer's authors put the underlying point bluntly: language models "struggle with basic functionality, such as arithmetic or factual lookup, where much simpler and smaller models excel."

Figure 5 — Parametric generation against a live system of record. The architectural distinction follows Lewis et al., NeurIPS 2020;³³ the tool-calling pattern follows Schick et al., NeurIPS 2023³⁶ and the augmented-language-model survey of Mialon et al., TMLR 2023.⁵³ The diagram is the authors' rendering of a hotel pricing exchange, not a figure from any cited paper.

Shuster and colleagues tested the effect directly, in a paper whose title is the finding: retrieval augmentation reduces hallucination in conversation. Their models "substantially reduce the well-known problem of knowledge hallucination in state-of-the-

art chatbots," as verified by human evaluation.³⁴ **PEER-REVIEWED** No percentage appears in that paper and none is quoted here.

Grounding is not a complete answer either. Patil and colleagues, presenting Gorilla at NeurIPS 2024, note

that even strong models show an "inability to generate accurate input arguments and their tendency to hallucinate the wrong usage of an API call."³⁷ Models can hallucinate the *call*, not only the prose. The Berkeley Function Calling Leaderboard's own conclusion, published at ICML 2025, is that

"while state-of-the-art LLMs excel at single-turn calls, memory, dynamic decision-making, and long-horizon reasoning remain open challenges."³⁸ A booking conversation is multi-turn and stateful. That is the hard case, not the easy one.

SECTION 07

Hallucination is not a bug you can patch out

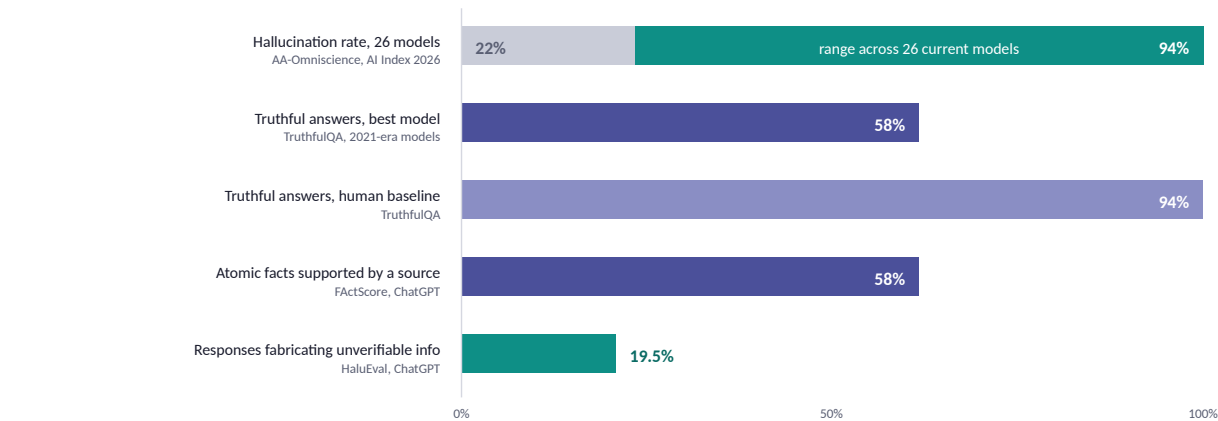
The common assumption is that hallucination was an early-model problem now largely solved. The 2026 measurements say otherwise, and a 2024 theoretical result says why.

The term has a precise definition. The standard survey, in *ACM Computing Surveys*, defines hallucination as "the generated content that is nonsensical or unfaithful to the provided source content," and splits it: **intrinsic** hallucination is output that contradicts the source, **extrinsic** is output that "cannot be verified from the source content."²⁵ **PEER-REVIEWED** Both categories are defined

relative to a source — which is the point. A system with no source of record can be neither faithful nor unfaithful to one. A later survey frames the consequence for deployment directly: models generate "plausible yet nonfactual content," raising "significant concerns over the reliability of LLMs in real-world information retrieval systems."²⁶

What the benchmarks measure

Each bar is a different benchmark measuring a different thing. They are not comparable to each other — read them as separate readings.



TruthfulQA and FActScore figures are from 2021–2023 model generations and are shown as historical readings, not current capability.

Figure 6 — Measured factual reliability. Hallucination range: Stanford HAI, *AI Index Report 2026*, Chapter 3, on the AA-Omniscience benchmark (6,000 questions across six domains).³² TruthfulQA: Lin, Hilton & Evans, ACL 2022.²⁹ FActScore: Min et al., EMNLP 2023.³¹ HaluEval: Li et al., EMNLP 2023.³⁰ These benchmarks measure different constructs on different model generations and must not be read as a single trend.

The AI Index figure is the one that matters for a 2026 buying decision: across **26 current models**, hallucination rates on AA-Omniscience range from **22 to 94 per cent**. The same chapter records that on a benchmark testing belief attribution, GPT-4o's accuracy falls from 98.2 per cent on true beliefs to 64.4 per cent on first-person false beliefs, and

DeepSeek R1 from over 90 per cent to 14.4 per cent.³² **INSTITUTIONAL**

THE FORMAL RESULT, AND ITS LIMIT

Kalai and Vempala, at the ACM Symposium on Theory of Computing in 2024, showed "an inherent statistical lower-bound on the rate that pretrained language models hallucinate certain types of facts, having nothing to do with the transformer LM architecture or data quality."²⁷ **PEER-REVIEWED**

The limit matters and we state it. The bound applies to facts that appear once or rarely in training data, and the authors are explicit that there is no statistical reason for hallucination on facts appearing repeatedly. So this does not mean models always hallucinate. It means they must hallucinate on exactly the class of fact a hotel cares about: a rate that exists in one database, for one property, on one night. A later paper from the same lead author argues the training and evaluation regime compounds this by rewarding "guessing over acknowledging uncertainty."²⁸

The counter-evidence, at full strength

Two findings cut against the case this report is making, and both belong in a buyer's decision.

Wüst and Bremser ran a preregistered experiment with 467 German participants on hotel booking support. Booking intention was **significantly lower with an AI chatbot than with a human agent** ($F(1,461) = 4.08, p = 0.044$), and fell more steeply for chatbots than humans when the booking situation turned negative.⁴⁵ **PEER-REVIEWED** Their recommendation is that complaints "should offer the option to be connected to a person." Germany is Türkiye's second-largest inbound market at 6.75 million visitors in 2025, so this is not a distant finding.⁵⁰

Longoni, Bonezzi and Morewedge, in the *Journal of Consumer Research*, identified the mechanism behind resistance to AI service as **uniqueness neglect** — a concern that AI cannot account for the consumer's particular circumstances — and found the resistance is eliminated when the AI is framed as personalised

or when it supports rather than replaces a human decision.⁴⁸ **PEER-REVIEWED**

Read together, these point somewhere specific. A system that reads the guest's actual dates, party composition and the property's actual policies is demonstrably accounting for their circumstances; and a system that hands complaints to a person addresses the case where chatbot acceptance is weakest. Neither finding argues against automation. Both argue against *undifferentiated* automation — and both are reflected in the specification in §11.

Two broader findings sit behind both. Dietvorst, Simmons and Massey established what they named **algorithm aversion**: people "are especially averse to algorithmic forecasters after seeing them perform, even when they see them outperform a human forecaster."⁴⁶ **PEER-REVIEWED** One visible error costs an automated system more trust than the same error costs a person — which raises the price of a single invented rate well above its face value. And a 2025 study of more than 48,000 respondents across 47

countries found that while 66 per cent of people use AI regularly, only **46 per cent are willing to trust it**.⁴⁹

INSTITUTIONAL

One further tension is worth surfacing. Wüst and Bremser recommend designing chatbots to be as human-like as possible. A 2025 study of 525 consumers found the opposite emphasis:

competence perception was the strongest driver of satisfaction ($\beta = 0.335$), while perceived social presence had essentially no effect ($\beta = 0.013$, $p = 0.765$).⁴⁷ The honest reading of both is that **getting the answer right beats seeming human** — which is, again, an argument about grounding rather than persona.

SECTION 08

The liability that arrived in 2024

Until February 2024 the question of who owns a chatbot's mistake was theoretical. A small-claims tribunal in British Columbia answered it, and the answer is an ordinary negligence standard.

Moffatt v. Air Canada, 2024 BCCRT 149 — the reasoning, in four steps

British Columbia Civil Resolution Tribunal, decision issued 14 February 2024.

1 The argument

Air Canada argued it could not be liable for what its chatbot told a customer — in effect, that the chatbot was a separate legal entity responsible for itself.

2 The response

"This is a remarkable submission."

"It makes no difference whether the information comes from a static page or a chatbot."

3 The standard

"the applicable standard of care requires a company to take reasonable care to ensure their representations are accurate and not misleading."

4 The finding

"I find Air Canada did not take reasonable care to ensure its chatbot was accurate."

Award: CAD \$812.02 in total.

WHAT IT MEANS OPERATIONALLY

A quote is a representation

If the bot names a price, the hotel has made a representation about that price. The channel does not change that.

"Reasonable care" is testable

Reading the price from the reservation system before sending it is evidence of care. Generating it is evidence of none.

Logs become the defence

An audit trail showing what the system returned, and what the guest was told, is what demonstrates the standard was met.

The award was trivial and this is a small-claims tribunal, not binding appellate precedent. Its significance is doctrinal: it settled that a chatbot has no separate legal existence from the business operating it, and that its accuracy is governed by the ordinary standard of care. Anyone citing it should say so.

Figure 7 — The reasoning in *Moffatt v. Air Canada*. Quotations are from the tribunal's decision, 2024 BCCRT 149, paragraphs 26–28.⁴¹ **OFFICIAL** The "what it means operationally" row is the authors' interpretation, not part of the decision.

The standard that is coming, and the one that already applies

From **2 August 2026**, Article 50 of the EU AI Act requires that AI systems intended to interact directly with natural persons be designed so that people "are informed that they are interacting with an AI system," unless that is obvious to a reasonably well-informed observer.⁴³ **OFFICIAL** Disclosure stops being a choice.

That matters commercially because disclosure has a measured cost. Luo, Tong, Fang and Qu, in *Marketing*

Science, exploited a field experiment across more than 6,200 customers and found undisclosed chatbots "as effective as proficient workers and four times more effective than inexperienced workers," while disclosure before the conversation "reduces purchase rates by more than 79.7 per cent" — because customers "perceive the disclosed bot as less knowledgeable and less empathetic."⁴⁴

PEER-REVIEWED

That study is from 2019, covers outbound sales calls in China, and predates current models. The 79.7 per cent figure will not replicate today and this report does not forecast it. What transfers is the *mechanism*: the penalty came from perceived lack of knowledge. Once disclosure is compulsory, the only remaining lever is for the disclosed system to be demonstrably knowledgeable — which is what grounding delivers and what an ungrounded assistant cannot.

For governance, the NIST AI Risk Management Framework organises practice into four functions: GOVERN, MAP, MEASURE and MANAGE, where MEASURE "employs quantitative, qualitative, or mixed-method tools, techniques, and methodologies

to analyze, assess, benchmark, and monitor AI risk and related impacts."⁴² **INSTITUTIONAL** MEASURE is the function an ungrounded assistant structurally cannot satisfy: you cannot benchmark the accuracy of an answer that has no source to check it against.

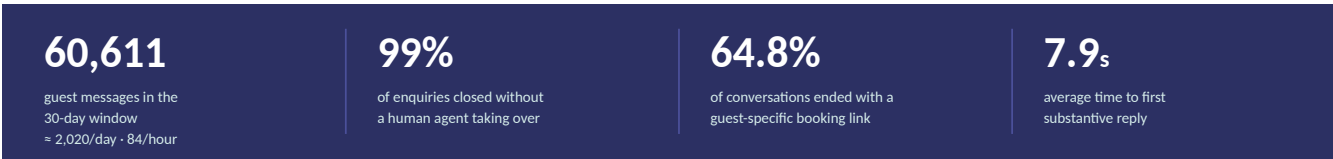
SECTION 09

The pilot: two hotels, seven months

What follows is the publisher's own operational data. It is presented with its measurement frame attached so that a reader can judge exactly what it supports — which is less than a vendor would like and more than nothing.

Pilot dataset — two hotel properties, approximately seven months of continuous operation

Headline metrics measured over one representative 30-day window within that period.



MEASUREMENT FRAME

Scope Two hotels under common ownership, not named at the operator's request. Channels: web chat, WhatsApp, Instagram.	Period ~7 months continuous live operation. Cumulative volume ≈ 424,000 messages, extrapolated from the measured window.	Grounding Every rate read from the hotels' live reservation system through an authenticated layer before being sent.
---	---	--

WHAT THIS DATASET DOES NOT ESTABLISH

X No control group — no causal claim about bookings or revenue	X Two properties, one ownership — not a sample, not generalisable
X No revenue attribution — no return-on-investment figure anywhere in this report	X No independent audit of the figures
X No guest-satisfaction instrument — nothing about experience quality	X One resort market — seasonality makes the extrapolation approximate

Figure 8 — The pilot dataset and its frame. **FIRST-PARTY** Publisher's own operational data. Definitions are given in full in §2. Read this as an existence proof that the architecture operates at volume with a low human-escalation rate — not as evidence of commercial return, which it cannot support.

Two observations are worth drawing out. First, the 7.9-second average sits against the 42-hour average response time measured across audited companies in §4 — a difference of roughly four orders of magnitude — and, more relevantly, against the 23 per cent of audited companies that never replied at all. The comparison that matters at three in the morning is not fast against slow. It is answered against unanswered.

Second, the 99 per cent autonomy figure should be read as what it is: the share of enquiries that did not require a human, in a deployment where the system was configured to hand over rather than guess. It is not a claim that one per cent of guests had a problem. Some share of that ninety-nine per cent were routine questions a printed card could have answered. The operationally interesting number is that the escalation path existed and was used.

The economics, on the conservative basis

Vendor return-on-investment arithmetic is normally computed on the cheapest plan and the lightest usage. This is computed on the most expensive plan and the heaviest usage, so that a reader's real cost is more likely to fall below it than above.

There are two cost lines and they behave differently. The **platform subscription** is fixed and published: the top tier, Enterprise Pro, is **\$2,990 a year** — an allowance of 6,000,000 conversations, 100 agent seats and 30 million AI training characters, with every feature included at every tier.⁵² The **AI usage cost** is variable, driven by prompt length, which

model is selected, and how many tool calls each conversation makes. The figure used below, **\$2,150 a month**, is the publisher's estimate computed for intensive use: long detailed prompts, a capable model, and heavy tool calling. It is an **upper bound**, not a typical bill.

The monthly cost of being available around the clock

Two ways of covering every hour. Currencies are shown as published and are not converted.

ONE PERSON ON DUTY, 24/7 — LABOUR COST

€9,391 – €15,215

per month, depending on market (\$3, Figure 2)

What it buys: **one person's attention**, one conversation at a time, in the languages that person speaks, minus leave, sickness and breaks.

Excludes recruitment, training and management overhead.

SYSTEM — TOP TIER, HEAVIEST USAGE ASSUMPTION

\$2,399 per month, upper bound

Platform subscription, Enterprise Pro	\$249
AI usage, intensive-use estimate	\$2,150

What it buys: every hour, every channel, every language, concurrent conversations without a queue.

COST PER GUEST MESSAGE, AT THE MEASURED VOLUME OF 60,611

\$0.040

per message, upper bound (platform + AI)

\$0.004

per message, platform subscription alone

€0.25

EU labour cost of 24/7 cover + same volume

The third box is not a like-for-like comparison: it divides the cost of one person's availability by a volume one person was never shown to be able to handle. See below.

Figure 9 — Monthly cost of round-the-clock availability. Platform pricing: publisher's published price list.⁵² AI usage: publisher's intensive-use estimate, **FIRST-PARTY**. Labour costs from Figure 2. Per-message figures are **AUTHOR'S CALCULATION** at the measured 60,611-message volume. No exchange-rate conversion is performed; the dollar and euro columns are not directly comparable and are shown together only to establish order of magnitude.

THE COMPARISON THIS FIGURE DOES AND DOES NOT MAKE

The two columns buy different things. The labour column buys **one person's availability**: one conversation at a time, in that person's languages. The system column buys availability plus the throughput actually observed — 2,020 messages a day across three channels.

Whether one person could have handled that throughput is precisely the question \$3 declines to answer, because no credible

independent benchmark exists to answer it with. So the honest statement is narrower than the arithmetic tempts: **at the same level of availability, the system cost sits roughly an order of magnitude below the labour cost, and it does not degrade when two guests write at once.** Anything stronger than that would require the headcount-equivalence claim this report has refused to make.

A deployment specification

Ten requirements, each traced to a finding in sections 3 to 10. Three of them exist because of evidence that cuts against AI guest service, not for it.

What a hotel should require of an AI guest communication system

Each requirement names the section and source it follows from. None names a vendor.

- 1 Every price read from the reservation system before it is sent**
A quote is a representation the hotel makes; generating one is evidence of no care taken — §6, §8
- 2 A clarifying question, never a guess, when the record does not cover the case**
A nightly rate is a fact appearing once; the statistical lower bound applies exactly there — §7, STOC 2024
- 3 The booking link encodes the same parameters as the quote**
Prevents the quote and the checkout diverging, which is where disputes originate — §6
- 4 A complete audit trail: what the system returned, and what the guest was told**
This is what demonstrates "reasonable care" after the fact — §8, Moffatt v. Air Canada
- 5 Rate and policy changes take effect without redeployment**
Parametric knowledge is stale by construction; rates change nightly — §6, Lewis et al. NeurIPS 2020
- 6 Multi-turn state held correctly across a whole booking conversation**
Single-turn calling is solved; multi-turn stateful interaction is the open problem — §6, BFCL ICML 2025
- 7 Coverage outside staffed hours, in the guest's own language**
The 60% comparison is against a 24-hour delay, and 23% of firms never replied at all — §3, §4
- 8 Human escalation always available, and automatic on complaints**
Booking intention falls further for bots than humans when the situation turns negative — §7, preregistered
- 9 AI identity disclosed to the guest**
EU AI Act Article 50 applies from 2 August 2026, and it is not optional — §8
- 10 Escalation rate, response latency and quote-to-record match reported as metrics**
NIST's MEASURE function requires benchmarking; an ungrounded system has nothing to benchmark against — §8

Requirements 1–6 are architecture. Requirement 7 is coverage. Requirements 8, 9 and 10 exist because of evidence against AI guest service.

A system that fails requirement 1 cannot be fixed by excelling at the other nine, because every other guarantee is downstream of where the price came from.
None of these requires a particular vendor. All ten can be put in a procurement document and tested in a trial.

Figure 10 — Evidence-derived deployment specification. Authors' synthesis of the findings in sections 3–10. Each requirement cites the section and underlying source it follows from.

Requirement 10 deserves a word. It is the one most often agreed in principle and never implemented, and it is the only one that tells you whether the other nine are working. Three numbers are enough: what proportion of conversations were escalated to a person, how long the first reply took, and — the one

nobody instruments — how often the price the guest was shown matched the price in the reservation system at that moment. A system that cannot report the third has not been grounded; it has been asserted to be.

Limitations and data gaps

Thirteen things this report cannot establish. Four of them concern the publisher's own data and three concern the publisher's own product category.

Limits of the pilot data

1. **No control group.** Nothing in §9 supports a causal claim about bookings, revenue or conversion, and

no return-on-investment figure appears anywhere in this report for that reason.

2. **Two properties under one ownership is not a sample.** No generalisation to other operators, segments, markets or property sizes is available from it.
3. **The 424,000 cumulative figure is an extrapolation** from one measured 30-day window

across a seven-month period, assuming a stable monthly run rate. In a seasonal resort market that assumption will not hold exactly, which is why the figure is rounded and described as approximate.

4. **The figures are not independently audited.** They are the publisher's own instrumentation.

Gaps in the published literature

5. **No peer-reviewed research links hotel enquiry response time to booking conversion.** The speed-to-lead evidence in §4 is cross-industry sales from 2011. Its transfer to hospitality is an assumption.
6. **No peer-reviewed study measures hallucination rates in hospitality chatbots specifically.** The nearest legitimate proxy is τ -bench, a customer-service agent benchmark covering retail and airline domains.³⁹ Any document implying hospitality-specific benchmark evidence exists is overstating the literature.
7. **No credible independent benchmark exists for chat throughput per agent.** The one figure in

circulation is vendor-derived (§3). One commercial benchmarking provider publishes a guest-services labour figure of 0.39 hours per occupied room, which is the closest hotel-specific productivity reference we located.¹⁸

VENDOR-REPORTED

8. **The central argument of this report is a synthesis, not a published finding.** Grounding reduces hallucination (§6), hallucination creates legal exposure (§8), and disclosure suppresses conversion unless the system is perceived as knowledgeable (§8) are three findings from three separate literatures. Assembling them into one chain is our argument. No one has published that chain as a result, and we do not present it as one.

Cautions on sources we do use

9. **Moffatt v. Air Canada is a small-claims tribunal decision**, not binding appellate precedent, and the award was CAD \$812.02. Its weight is doctrinal, not financial.
10. **The 79.7 per cent disclosure penalty is from 2019**, covers outbound sales calls in China, and predates current models. The mechanism transfers; the magnitude will not.⁴⁴
11. **EU AI Act Article 50 does not apply until 2 August 2026**, so no evidence of how mandatory disclosure actually changes guest behaviour can exist yet. Anyone offering you such evidence today is offering you a forecast.
12. **HOTREC's European shortage figures are supplied by national member associations**, not by a harmonised statistical survey, and cross-country comparability is weak.¹⁵ The Eurostat vacancy series is the harmonised alternative and is used wherever possible. Eurostat's minimum-wage series also annualises 14-payment countries

over twelve months, so its figures differ from national headline wages by design.^{8,9}

13. **Older benchmark figures are labelled as such.** TruthfulQA and FActScore results in Figure 6 are from 2021–2023 model generations. Only the AI Index range covers current models.

THE FINDING THAT SHOULD TEMPER ANY AUTONOMY CLAIM

τ -bench evaluates language agents against real API tools and written policy documents in customer-service domains, scoring by whether the final database state matches an annotated goal. Its result: "even state-of-the-art function calling agents (like gpt-4o) succeed on <50% of the tasks, and are quite inconsistent (pass@8 <25% in retail)".³⁹

PEER-REVIEWED WebArena reports a comparable picture for open web tasks — a best agent success rate of 14.41 per cent against human performance of 78.24 per cent.⁴⁰ Both are older model generations, and both make the same point: **reliability, not capability, is the product**. An agent that succeeds once is not an agent that succeeds every time, and a hotel is buying every time. This is the strongest argument for narrowing the scope of what an AI system is permitted to do — and for requirement 2 in §11.

SECTION 13

References

All sources accessed and verified September 2026. Peer-reviewed work is cited to the venue that published it; where a finding reached us through a university news office or trade outlet, both are named.

- 1 Oldroyd, J.B., McElheran, K. & Elkington, D. (2011), "The Short Life of Online Sales Leads", *Harvard Business Review* 89(3), p. 28. hbr.org/2011/03/the-short-life-of-online-sales-leads
- 2 BYU ScholarsArchive, faculty publication record for the above. scholarsarchive.byu.edu/facpub/9711
- 3 Elkington, D. & Oldroyd, J., Lead Response Management Study, InsideSales.com. **Vendor platform data**, not an MIT publication. insidesales.com/lead-response-management-study
- 4 InsideSales.com / XANT, *Lead Response Report 2014* (2013 data; 9,538 companies). **Vendor-reported**. resources.insidesales.com/2014-Lead-Response-Report.pdf
- 5 Eurostat, Hourly labour costs (Statistics Explained), 2025. ec.europa.eu/eurostat/Hourly_labour_costs
- 6 Eurostat, dataset `lci_lci_level`, NACE Rev. 2 Section I (accommodation and food service activities), 2025. ec.europa.eu/eurostat/databrowser/view/lci_lci_level
- 7 T.C. Çalışma ve Sosyal Güvenlik Bakanlığı, 2026 asgari ücret ve işverene maliyeti. csgeb.gov.tr/asgari-ucret-2026.pdf
- 8 Eurostat, Minimum wage statistics (Statistics Explained), 1 July 2026. ec.europa.eu/eurostat/Minimum_wage_statistics
- 9 Eurostat, dataset `earn_mw_cur`, monthly minimum wages, 2026-S2. ec.europa.eu/eurostat/databrowser/view/earn_mw_cur
- 10 Eurostat, Job vacancy statistics, Q1 2026. ec.europa.eu/eurostat/Job_vacancy_statistics
- 11 US Bureau of Labor Statistics, JOLTS Table 4 (quits), July 2026 preliminary. bls.gov/news.release/jolts.t04.htm
- 12 US Bureau of Labor Statistics, JOLTS Table 16 (annual average job openings rates, 2021–2025). bls.gov/news.release/jolts.t16.htm
- 13 American Hotel & Lodging Association, Front Desk Feedback survey (n=246 hoteliers, late February 2026), published 17 March 2026. ahla.com/news/rising-cost-staffing-challenges-persist
- 14 American Hotel & Lodging Association, staffing shortage survey (n=282, Dec 2024–Jan 2025), conducted with a recruiting-software partner. ahla.com/news/65-surveyed-hotels-report-staffing-shortages
- 15 HOTREC Hospitality Europe, *Skills and Labour Shortages: A Roadmap for Action*, January 2026. Member-supplied estimates. hotrec.eu/skills-and-labour-shortages-roadmap-2026.pdf
- 16 Gans, N., Koole, G. & Mandelbaum, A. (2003), "Telephone Call Centers: Tutorial, Review, and Research Prospects", *Manufacturing & Service Operations Management* 5(2), 79–141. DOI 10.1287/msom.5.2.79.16071. pubsonline.informs.org/doi/10.1287/msom.5.2.79.16071

- 17 **Call Centre Helper**, live chat benchmarks. Figures attributed by the publication to commercial vendors. **Vendor-derived**. callcentrehelper.com/benchmarks-improve-live-chat-metrics
- 18 **HotelData.com**, US hotel labour benchmarks, Jan–Sep 2025 (~5,000 US hotels). **Commercial benchmarking provider**. hoteldata.com/benchmark/us-hotel-labor-benchmarks-statistics
- 19 **DataReportal** (Kepios / We Are Social / Meltwater), global social media users. datareportal.com/social-media-users
- 20 **DataReportal**, *Digital 2026: Turkey*. datareportal.com/reports/digital-2026-turkey
- 21 **Meta Platforms**, Q1 2026 earnings call transcript, 29 April 2026. investor.atmeta.com/META-Q1-2026-Earnings-Call-Transcript.pdf
- 22 **Oracle Hospitality with Skift**, *Hospitality in 2025* (5,266 consumers, 633 hotel executives, nine markets). **Vendor-reported**. prnewswire.com/travelers-want-high-tech-low-touch-hotel-stays
- 23 **Hospitality Technology**, Lodging Technology Study (respondents representing 10,000+ properties). hospitalitytech.com/2024-lodging-tech-study
- 24 **Wyndham Hotels & Resorts**, Owner Trends Report, January 2026, via Hotel Technology News. Sample size not disclosed. **Vendor-reported**. hoteltechnologynews.com/98-of-hotels-have-begun-using-ai
- 25 **Ji, Z. et al.** (2023), "Survey of Hallucination in Natural Language Generation", *ACM Computing Surveys* 55(12), Article 248. DOI 10.1145/3571730. arxiv.org/abs/2202.03629
- 26 **Huang, L. et al.**, "A Survey on Hallucination in Large Language Models", *ACM Transactions on Information Systems*, 2025. DOI 10.1145/3703155. arxiv.org/abs/2311.05232
- 27 **Kalai, A.T. & Vempala, S.S.** (2024), "Calibrated Language Models Must Hallucinate", *Proceedings of the 56th Annual ACM Symposium on Theory of Computing (STOC)*. arxiv.org/abs/2311.14648
- 28 **Kalai, A.T., Nachum, O., Vempala, S.S. & Zhang, E.** (2025), "Why Language Models Hallucinate". Preprint; no peer-reviewed venue. arxiv.org/abs/2509.04664
- 29 **Lin, S., Hilton, J. & Evans, O.** (2022), "TruthfulQA: Measuring How Models Mimic Human Falsehoods", *ACL 2022*. arxiv.org/abs/2109.07958
- 30 **Li, J. et al.** (2023), "HaluEval: A Large-Scale Hallucination Evaluation Benchmark for Large Language Models", *EMNLP 2023*. arxiv.org/abs/2305.11747
- 31 **Min, S. et al.** (2023), "FactScore: Fine-grained Atomic Evaluation of Factual Precision in Long Form Text Generation", *EMNLP 2023*. arxiv.org/abs/2305.14251
- 32 **Stanford Institute for Human-Centered AI**, *Artificial Intelligence Index Report 2026*, Chapter 3: Responsible AI. hai.stanford.edu/ai_index_report_2026_chapter_3_responsible_ai.pdf
- 33 **Lewis, P. et al.** (2020), "Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks", *NeurIPS 2020*. arxiv.org/abs/2005.11401
- 34 **Shuster, K., Poff, S., Chen, M., Kiela, D. & Weston, J.** (2021), "Retrieval Augmentation Reduces Hallucination in Conversation", *Findings of EMNLP 2021*, 3784–3803. DOI 10.18653/v1/2021.findings-emnlp.320. aclanthology.org/2021.findings-emnlp.320
- 35 **Gao, Y. et al.**, "Retrieval-Augmented Generation for Large Language Models: A Survey". Preprint. arxiv.org/abs/2312.10997
- 36 **Schick, T. et al.** (2023), "Toolformer: Language Models Can Teach Themselves to Use Tools", *NeurIPS 2023*. arxiv.org/abs/2302.04761
- 37 **Patil, S.G., Zhang, T., Wang, X. & Gonzalez, J.E.** (2024), "Gorilla: Large Language Model Connected with Massive APIs", *NeurIPS 2024*. arxiv.org/abs/2305.15334
- 38 **Patil, S.G. et al.** (2025), "The Berkeley Function Calling Leaderboard: From Tool Use to Agentic Evaluation of Large Language Models", *ICML 2025*, PMLR 267, 48371–48392. proceedings.mlr.press/v267/patil25a.html
- 39 **Yao, S., Shinn, N., Razavi, P. & Narasimhan, K.** (2025), "t-bench: A Benchmark for Tool-Agent-User Interaction in Real-World Domains", *ICLR 2025*. arxiv.org/abs/2406.12045
- 40 **Zhou, S. et al.** (2024), "WebArena: A Realistic Web Environment for Building Autonomous Agents", *ICLR 2024*. arxiv.org/abs/2307.13854
- 41 **Moffatt v. Air Canada**, 2024 BCCRT 149, British Columbia Civil Resolution Tribunal, 14 February 2024, File SC-2023-005609. canlii.org/2024bccrt149
- 42 **NIST**, *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, NIST AI 100-1, January 2023. DOI 10.6028/NIST.AI.100-1. nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf
- 43 **European Union**, Regulation (EU) 2024/1689 (Artificial Intelligence Act), Article 50. Obligations apply from 2 August 2026. eur-lex.europa.eu/eli/reg/2024/1689/oj/eng
- 44 **Luo, X., Tong, S., Fang, Z. & Qu, Z.** (2019), "Frontiers: Machines vs. Humans: The Impact of Artificial Intelligence Chatbot Disclosure on Customer Purchases", *Marketing Science* 38(6), 937–947. DOI 10.1287/mksc.2019.1192. pubsonline.informs.org/doi/10.1287/mksc.2019.1192
- 45 **Wüst, K. & Bremser, K.** (2025), "Artificial Intelligence in Tourism Through Chatbot Support in the Booking Process—An Experimental Investigation", *Tourism and Hospitality* 6(1), 36. DOI 10.3390/tourhosp6010036. mdpi.com/2673-5768/6/1/36
- 46 **Dietvorst, B.J., Simmons, J.P. & Massey, C.** (2015), "Algorithm Aversion: People Erroneously Avoid Algorithms After Seeing Them Err", *Journal of Experimental Psychology: General* 144(1), 114–126. DOI 10.1037/xge0000033. pubmed.ncbi.nlm.nih.gov/25401381
- 47 **Al-Oraini, B.S.** (2025), "Chatbot dynamics: trust, social presence and customer satisfaction in AI-driven services", *Journal of Innovative Digital Transformation* 2(2), 109. DOI 10.1108/JIDT-08-2024-0022. emerald.com/jidt/chatbot-dynamics
- 48 **Longoni, C., Bonezzi, A. & Morewedge, C.K.** (2019), "Resistance to Medical Artificial Intelligence", *Journal of Consumer Research* 46(4), 629–650. DOI 10.1093/jcr/ucz013. academic.oup.com/jcr/46/4/629
- 49 **Gillespie, N., Lockey, S. et al.** (2025), *Trust, attitudes and use of artificial intelligence: A global study 2025*, University of Melbourne with KPMG. 48,000+ respondents, 47 countries. kpmg.com/trust-in-artificial-intelligence
- 50 **T.C. Kültür ve Turizm Bakanlığı with TÜİK**, 2025 border statistics, released 30 January 2026; accessed via secondary reproduction. turizmguzel.com/2025-yilinda-turkiye-gelen-turist-sayisi
- 51 **Instituto Nacional de Estadística (Spain)**, FRONTUR, full year 2025 (provisional). ine.es/dyns/Prensa/FRONTUR1225.htm
- 52 **Vera**, published pricing. vera.support/pricing
- 53 **Mialon, G. et al.** (2023), "Augmented Language Models: a Survey", *Transactions on Machine Learning Research*. arxiv.org/abs/2302.07842

Implementation resources

The resources below are the publisher's own material, listed against the section and specification requirement each relates to. They are operational and product documentation, not research sources; nothing in sections 3 to 12 is cited to any of them.

THE SYSTEM DESCRIBED IN THIS REPORT — §6, §9

AI assistant: grounding, escalation and language handling
vera.support/ai-chatbot/

Hotels and hospitality
vera.support/solutions/hotels/

Full feature list
vera.support/features/

COVERAGE AND CHANNELS — §3, §5, REQUIREMENT 7

Channel integrations: WhatsApp, Instagram, Messenger, Telegram and others
vera.support/apps/

Travel agencies and tour operators
vera.support/solutions/travel/

Restaurants and food service within the property
vera.support/solutions/restaurants/

THE PRICING USED IN §10

Published price list — every tier includes every feature
vera.support/pricing/

Platform overview
vera.support

ADJACENT DEPLOYMENTS — CONTEXT ONLY

Clinics and healthcare
vera.support/solutions/clinics/

Events and venues
vera.support/solutions/events/

All industry solutions
vera.support/solutions/

REUSE, CITATION AND CORRECTIONS

Quotation, excerpting and reproduction are permitted with attribution. Suggested citation: *The Grounded Front Desk: Guest Communication, Response Latency and Factual Reliability in Hotel Operations*, 2026. Vera, September 2026. Corrections are welcome and will be published — if you can supply a primary source for any item in §12, particularly peer-reviewed research linking hotel enquiry response time to booking conversion, we will update this report and credit the correction.