

Trusted Answers at Enterprise Scale

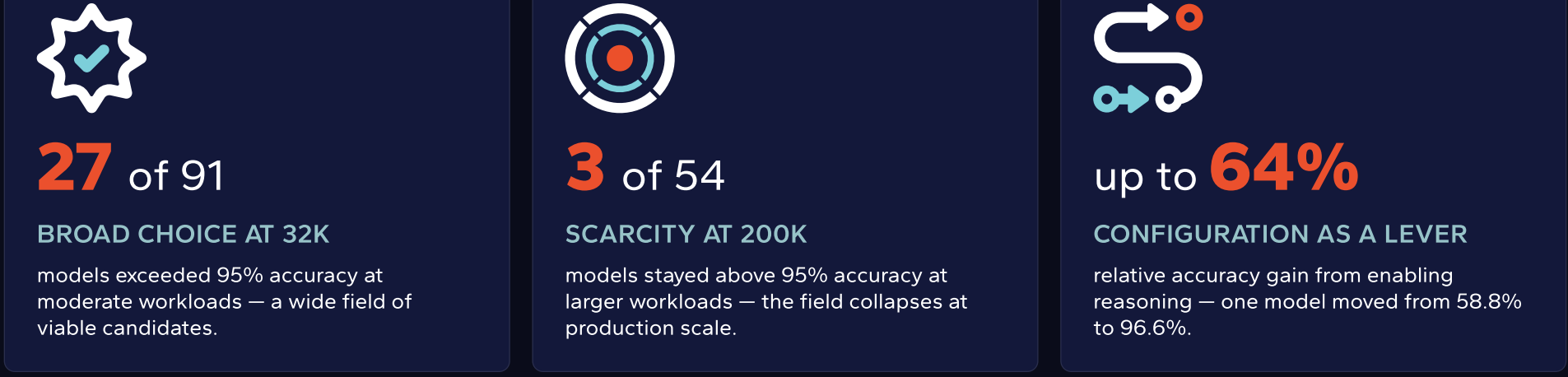
Retrieval reliability for enterprise AI, measured on a Dell and Broadcom stack

Enterprise AI's most common job is answering questions from the organization's own documents: contracts, policies, operational reports, and records. The business case rests on one question — **can the answers be trusted?**



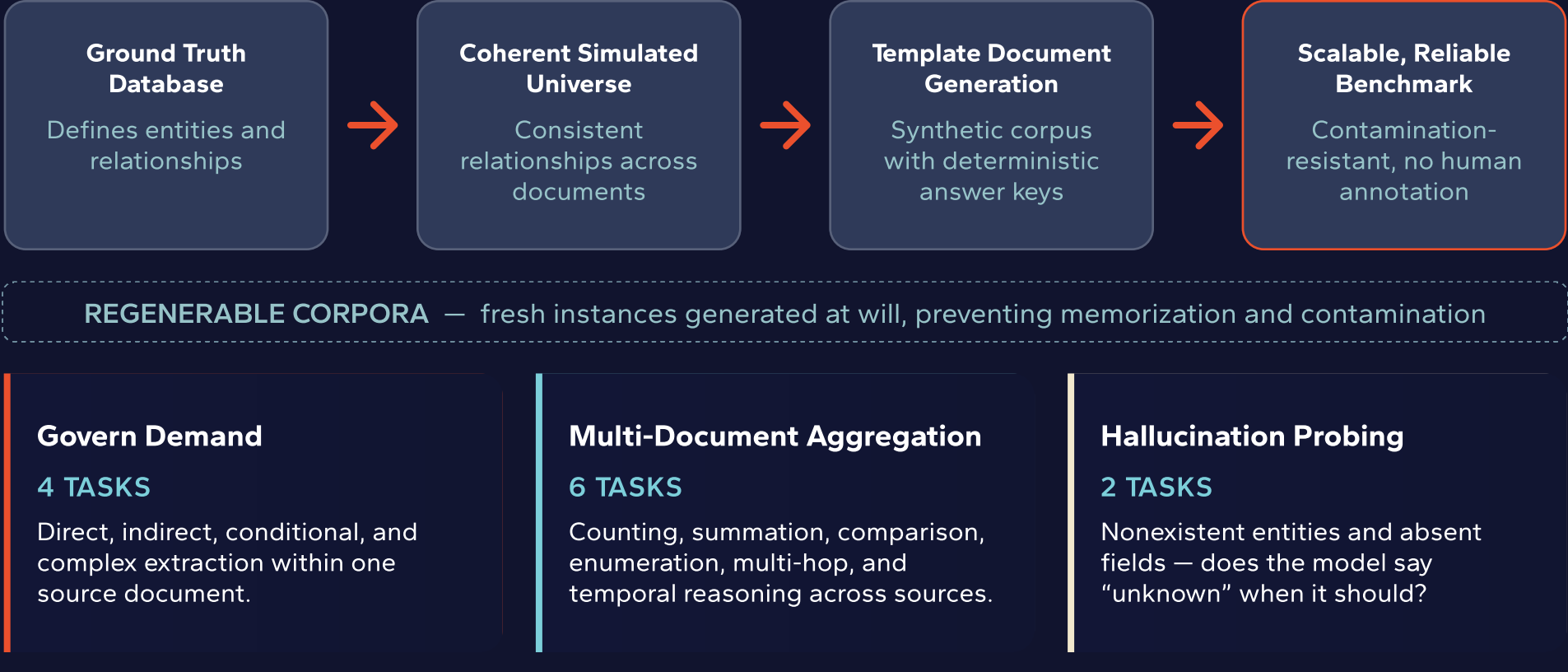
THE RETRIEVAL TRUST GAP

Retrieval quality holds up at moderate document scale and thins out sharply as corpora grow. Deployment configuration moves the number as much as model choice does.



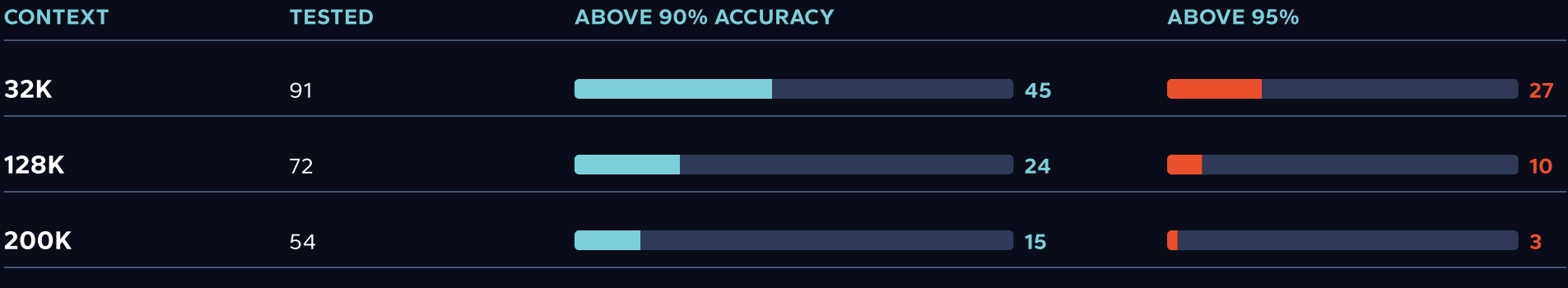
HOW RIKER WORKS

Traditional benchmarks start from existing documents, then annotate answers by hand — slow, hard to scale, and vulnerable to training-data contamination. RIKER inverts the process: it generates documents from predefined ground-truth answers, so grading stays deterministic and every run is fresh.



THE FIELD NARROWS AS DOCUMENTS GROW

Testing covered 91 models at 32K context, 72 at 128K, and 54 at 200K, using corpora of commercial leases, HR records, and facility field reports. Advertised context windows did not predict which models held their accuracy.



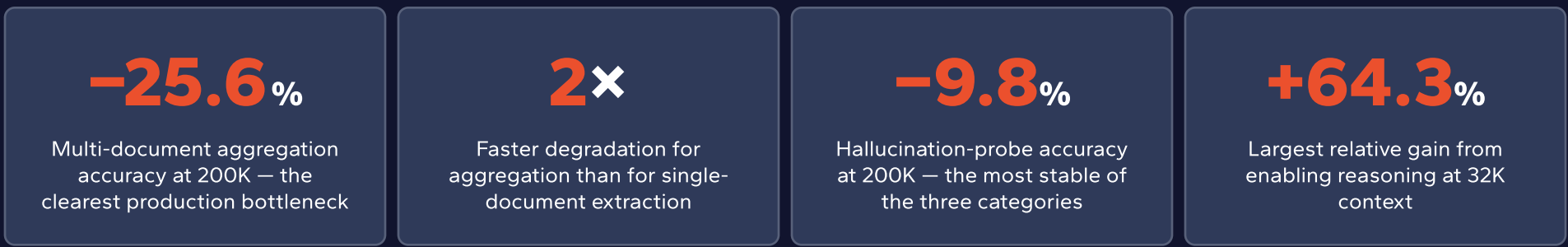
At 200K context, only **three of 54 models** sustained the 95% standard. A shortlist built on 32K results alone will not survive contact with production-scale corpora.

OVERALL MEAN ACCURACY — 200K CONTEXT



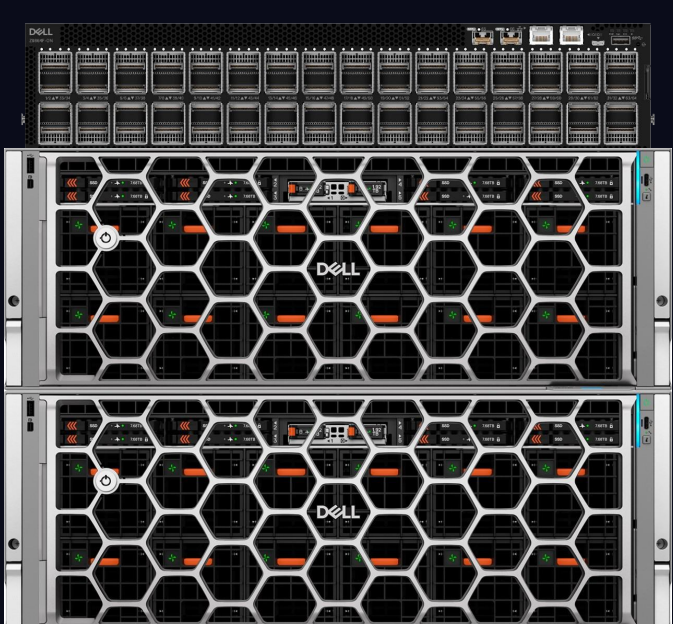
WHERE ACCURACY BREAKS DOWN

Averaged across all evaluated models, relative to 32K performance. Cross-document work degrades fastest; knowing when an answer is absent proves the most durable capability.



THE TESTED OPEN-MODEL STACK

Open models ran in the Signal65 AI Lab on an on-premises stack, keeping sensitive documents under the organization's control. Proprietary models were run through their native API endpoints.



DELL POWERSWITCH Z9864F-ON + 2x DELL POWEREDGE XE7745

Dell PowerSwitch Z9864F-ON — The Fabric

Built on Broadcom Tomahawk 5 silicon and running Enterprise SONiC 4.4.0, the switch interconnects the cluster over 400GbE RoCEv2. Long-context inference raises memory demand for the KV cache, and distributed serving adds inter-node communication — the fabric is what keeps that traffic predictable.

Dell PowerEdge XE7745 — The Compute

A high-memory, multi-GPU platform for large-context inference and high-volume retrieval. Testing also used PowerEdge XE9680 with NVIDIA H200 and AMD MI300X. Each XE7745 node in the tested cluster:

- Dual AMD EPYC 9555 · 2.3 TB DDR5 memory
- 8× NVIDIA RTX Pro 6000 Blackwell Server
- 8× Broadcom BCM57608 400GbE RoCEv2 network controllers
- Ubuntu 24.04 LTS · CUDA 13.0 / Driver 580.95.05 · vLLM v0.10.2

FROM RISK TO RESOLUTION

The findings map directly onto the exposures an enterprise retrieval program places on the budget owner's desk.

Business Risk	What the Testing Shows
Procurement Risk	Advertised context windows are capacity limits, not evidence of retrieval quality. Specify accuracy at corpus scale in evaluation criteria.
Pilot-to-Production Risk	Single-document pilots overstate readiness; cross-document aggregation fell 25.6% on average at 200K context.
Agentic Program Risk	Per-step errors compound: five retrieval-dependent steps at 90% reliability complete cleanly only about 59% of the time.
Compliance and Audit Risk	Recognizing absent information is a separate capability from retrieval, and should be measured on its own.

Validate Retrieval Before Delegating Autonomy

Agents act on what they retrieve, so a retrieval error becomes an action error. Evaluate models, configurations, and infrastructure together against the workloads the organization actually intends to run — not against specifications alone.