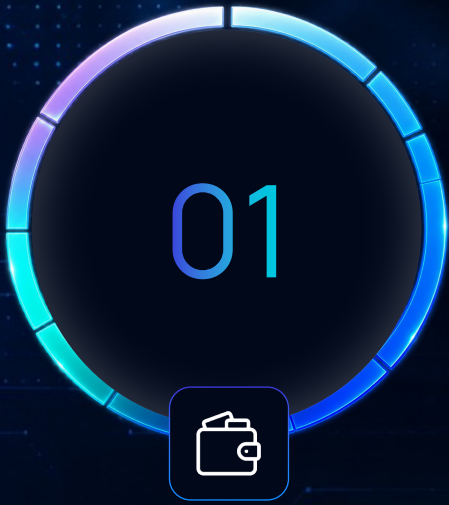


HEXAWARE

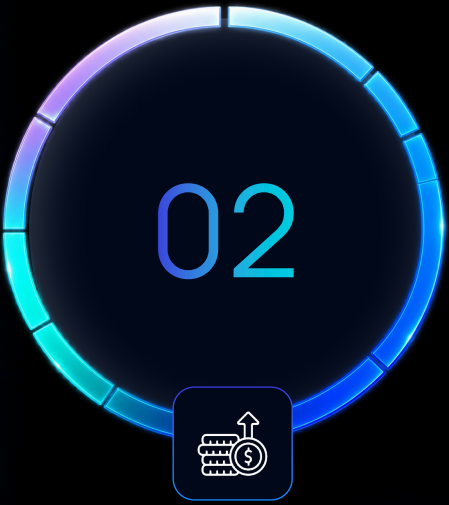
Welcome!

August 2026

Our AI Strategy is Designed to Achieve Two Outcomes



Create and sustain
new revenue lanes
from AI



Improve
profitability



Hexaware AI Strategy for Customers

AI for IT

Guardrails First, Autonomy Next



Customers define the guardrails, governance, and priorities; Hexaware integrates AI into every layer of IT operations.

AI for Business

Building Agentic AI Models Together



Customers prioritize agentic AI models that drive competitive differentiation; Hexaware partners to design and build them.

Foundational Layer for AI

Underpinning these pillars is Infinite Trust, a foundational layer creating the conditions for AI to operate securely, responsibly, and at enterprise scale.



Data Readiness
for AI

Security

Observability &
Governance



700+
Team Strength



27
Average Age

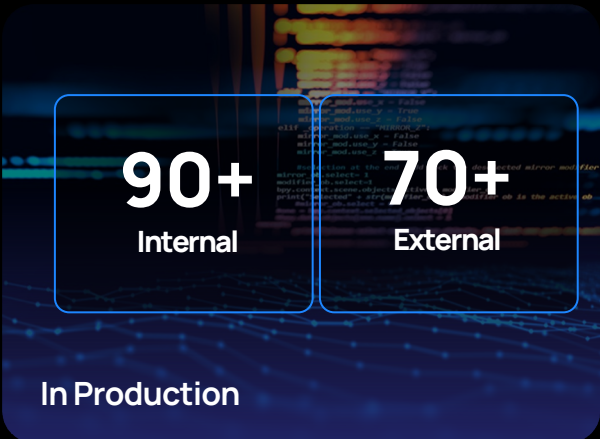


Best in Class Infra

OUR LAB IS THE ENGINE BEHIND
OUR **AGILITY**

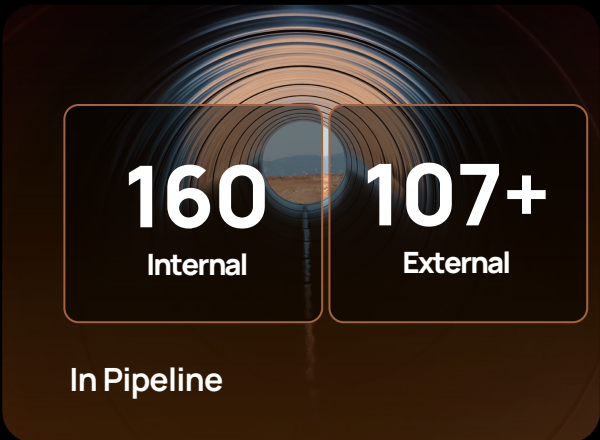


Every service we
bring to market has
its origins in the
AI Lab.



90+ Internal	70+ External
------------------------	------------------------

In Production



160 Internal	107+ External
------------------------	-------------------------

In Pipeline

ANTHROPIC			UNITREE
 FACTORY	 Microsoft	Google	 NVIDIA

Zero Friction, Infinite Momentum.

HEXAWARE



Zero Friction, Infinite Momentum.

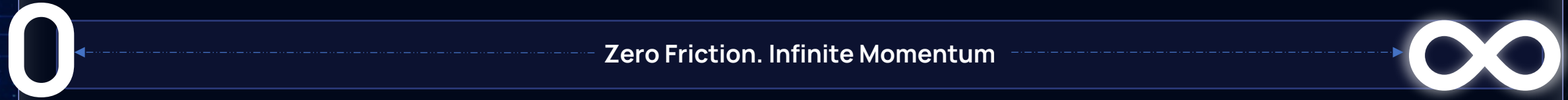
HEXAWARE

● New revenue stream ● Deflationary with volume offset ● Deflationary

AI for IT

Guardrails First, Autonomy Next

Customers define the guardrails, governance, and priorities; Hexaware integrates AI into every layer of IT operations.



Zero Tech Debt
TAM ~\$65B

Zero Vulnerability
TAM ~\$25B

Zero Backlog

Zero Defects

Zero Tickets

Zero License
TAM ~\$48B

Infinite Trust

Data Readiness for AI
TAM ~\$40B

Observability and Governance
TAM ~\$60B

Security
TAM ~\$80B

Zeroivity™— A Single Context, Delivers Multiple Services

HEXAWARE

● New revenue stream ● Deflationary with volume offset ● Deflationary

ZEROVITY - AI delivery layer for governed coordination across the estate

Code
Comprehension

Knowledge Graph

AI / LLM Harness

Agent Studio

Blueprint Co-Pilot

Parity & Validation

PRODUCT MODELS · Variants Built On One Foundation

RapidX®

Legacy Modernization

Amaze® Cloud

Non-Functional
Modernization

Code Forensics

Code Intelligence

Tensai®

ReasoningOps

Zeroivity Starter

Assessment Only

GTM PATTERNS · Offerings Built on the Foundation and Products

Zero Tech Debt

Zero Vulnerability

Zero Backlog

Zero Defects

Zero Tickets

Zero License

AI for Business: Solving the Hard Problems. Entering New Markets.

Solving the hard problem

Fab yield is the single biggest lever on revenue—yet throughput is capped, know-how sits in silos, and a 157k skilled-worker shortfall is anticipated by 2030. Agentic factory ops catch issues at the machine before they become incidents.

4–6 hr. → <20 min

Fab MTTR; 30–40% engineer capacity reclaimed

Entering new markets

A \$50B market-data industry growing ~6.5% a year, where clients pay ~\$100M per \$1T AUM with no view of contracts, usage, or attribution. We start with MacroIQ spend optimization, then become a provider: AI-native datasets and proprietary indices.

\$50B TAM

20–50% cost out, 70% faster to market

Tokenomics Determines IT Budgets

Dual Imperative for Hexaware: Participate in tokenomics and provide value to customer

01



Participate in tokenomics
 Alternate pricing models

02

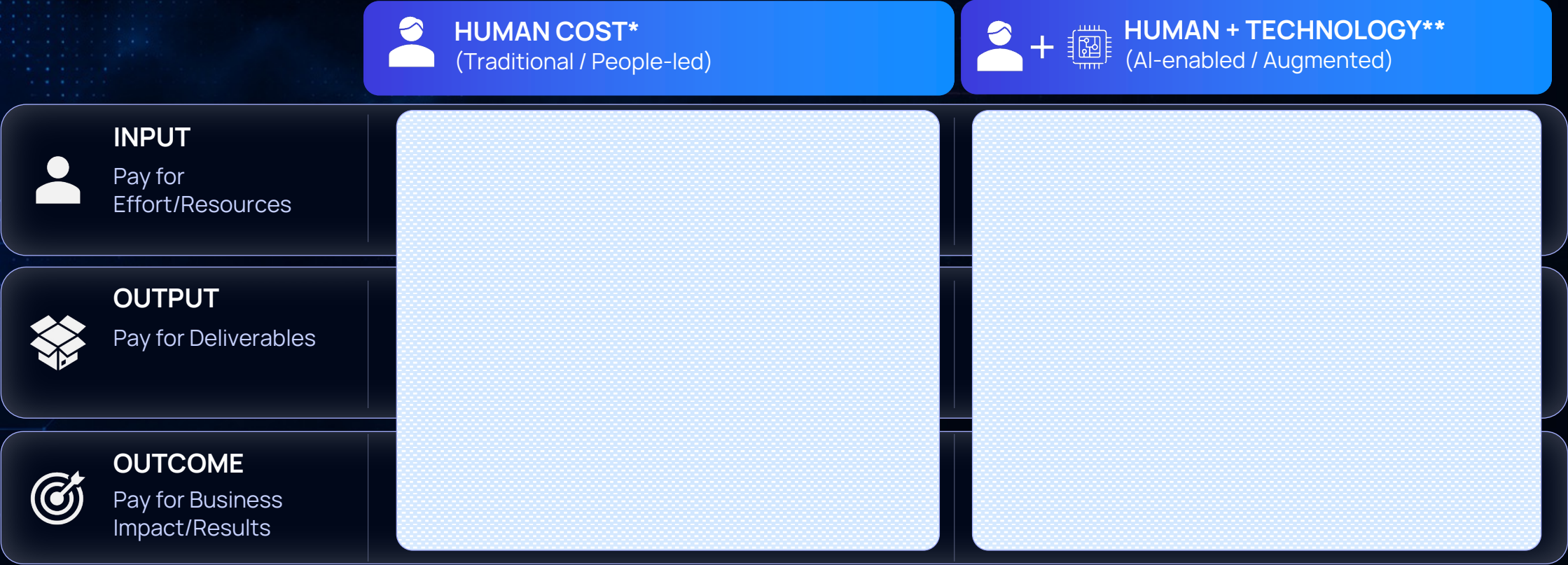


Provide value to customer
 Value priced into every token



1. Participate in Tokenomics Through Our Pricing Model

AI is shifting commercial models from paying for effort to paying for measurable business outcomes



Traditional Pricing
Effort based Model



Modern Pricing
Output-based Model



AI-native Pricing
Value-based models

2. Provide Value Through Tokenomics

Three levers set the unit economics of every AI engagement

1

Design for reduced token usage

Prompt, context, and retrieval engineered to minimize tokens per transaction—the single largest driver of run-rate cost.

2

2

Leverage Harness to drive economics

Retrieval, orchestration, caching, routing, and guardrails—tuning the harness moves cost more than swapping the model.

3

Deploy the right language model

Route most steps to small and micro language models; reserve frontier LLMs for the few that need them.



Token Efficiency Techniques We Apply [e.g., On Major Models]

Token spend is a managed engineering variable: measured per agent, optimized by technique, and governed by budget gates—on Claude Code via Vertex AI.

THE ECONOMIC MODEL

**Metered per agent**

Harness cost management tracks token spend by agent, task and phase.

**Budgeted, not unbounded**

Each workflow has a token budget; gates warn/halt before overrun.

**Value-indexed**

Spend is reported against outcomes (PRs, defects avoided)—cost per unit of value.

**Transparent weekly**

See spend, trend and ROI from Week 1—no surprises.

TOKEN-EFFICIENCY TECHNIQUES WE APPLY

Rapdix Agents/ Harness Capability

Claude Code Capability

**Context compaction**

Strategic-compact summarizes long sessions so context stays lean, not bloated

**Targeted file navigation**

Claude code greps/globs to open only the files and symbols a task touches—no whole repo loads, no RAG infra

**Prompt + result caching**

Vertex AI prompt caching reuses stable context (CLAUDE.md, schemas) across calls

**Right-sized model routing**

Cheaper models for boilerplate/lint; Claude reserved for reasoning-heavy work

**Spec-bounded tasks**

Small specs cap scope per call—fewer tokens, fewer retries, less rework

**Reusable agent blocks**

Pre-built LM agent blocks avoid re-deriving context every run

**Batch + parallelize**

Group related changes; batch reviews to amortize fixed context cost

**Diff-only operations**

Operate on diffs/patches rather than re-emitting whole files

**Enterprise context inheritance**

BUD topology, EDMX, and OpenBox chain pre-loaded once, not re-explained per task

It Is Not Just a Tech Talent Transformation

“ Leadership

10x Velocity,
Experimenting,
Risk-taking,
Learning

“ Client Teams

AI Thought Leaders,
Proactive, Use AI in
Everything

“ New Skills

250 FDEs and
1000 AI
Architects in
12 Months

“ Everyone

Change Language,
Highest AI Literacy



~ 50% revenues
from AI or infused
with AI

Introducing

The Zero Friction portfolio and the leaders driving it

Zerovity™

Siddharth Dhar



Zero License

Vaibhav Bhatnagar



Zero Tech Debt & Zero Backlog

Sanjay Salunkhe



Zero Tickets

Girish Ravindran



Infinite Trust

Girish Pai



AI in Business

Shantanu Baruah, Eravi Gopan,
Ravi Vaidyanathan, and Kush Gupta



Innovation

Satyajith M



- New revenue stream
- Deflationary, volume offset
- Deflationary