

J4

Logic Field Processor

Powering brain-scale intelligence

Dedicated silicon for J4 Stochastic Learning
Massive computation at the edge
Lower power and cooling burden in data centers

\$150M financing request

Patent-pending learning
architecture → commercial
silicon

THE CONSTRAINT HAS MOVED

Generative AI has hit a physical limit

The next AI bottleneck is no longer only the algorithm. It is the infrastructure required to run it.

950 TWh

Projected global data-center
electricity use by 2030

IEA

11.8%

Estimated share of U.S. electricity
used by data centers by 2030

Berkeley Lab

\$487B

Projected 2026 AI infrastructure
spending

IDC

**Same workload
More infrastructure**

+

**Different learning method
New performance curve**

AI is consuming the operating envelope

What can be powered, cooled, connected, permitted, and built now limits what can be computed.

2×

Global data-center electricity
by 2030
485 → 950 TWh
IEA

9.5–15.3%

U.S. electricity scenario
range by 2030
Berkeley Lab

33–40%

Facility energy may be used
for cooling
ARPA-E

**POWER
GRID**

**WATER
COOLING**

The cost of AI is no longer just compute.

It is the infrastructure required to make compute possible.

The cloud is too far from the edge

Machines, sensors, vehicles, devices, and facilities need decisions where their data is created.

CLOUD AI

Centralized
Connected
Distant

Network round-trip latency
Bandwidth and transfer cost
Privacy and data-residency exposure
Loss of autonomy when disconnected

EDGE AI

Local
Autonomous
Real-time

J4 objective
Massive local computation
inside practical power limits

Today's accelerators optimize the same paradigm

GPUs, TPUs, NPUs, and cloud ASICs are extraordinary—yet they largely accelerate the same dominant workload.

DOMINANT WORKLOAD

- Matrix and tensor operations
- Deep-learning frameworks
- High-bandwidth memory movement
- Centralized scale-out processing

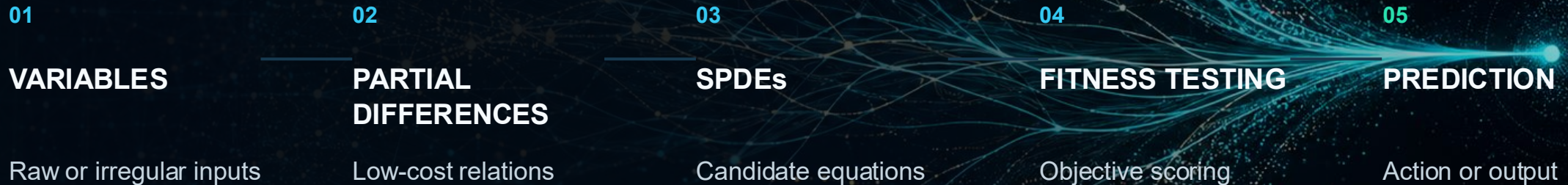
J4 ALTERNATIVE

- Stochastic partial differences
- Relationship discovery fabric
- Fitness-based pruning
- Edge-to-data-center processing

LFP is designed for a different learning architecture—not merely to run generative AI faster.

J4 Stochastic Learning changes the boundaries

Rather than relying only on repeated tensor operations, J4 searches for predictive stochastic relationships among variables.



Patent-disclosed advantages

Reduced preprocessing • reduced feature engineering • optional unsupervised operation • low-cost difference operations • inherent parallelism

LFP turns Stochastic Learning into silicon

A proposed compute fabric specialized for relationship generation, objective scoring, and pruning.

Ingest + local memory

Partial-difference engines

Relationship fabric

Fitness processors

Pruning + state memory

Prediction / control output

PLATFORM SCOPE

Chip architecture • compiler • runtime • profiler • SDK • reference systems

The round funds one decisive performance question

Can J4 independently reproduce a step-function gain in useful predictions per joule at equivalent task quality?

TARGET TO VALIDATE

1,000x+

Useful predictions per joule
on selected workloads

01 Independent validation

Energy + task quality

02 FPGA scale-out

Kernel parallelism

03 First silicon

Thermal + system proof

The proof comes before the largest silicon commitments.

A PLATFORM, NOT A SINGLE CHIP

Customers can evaluate, program, deploy, and scale J4

The product family creates a path from software evaluation to edge and data-center deployment.

01

SDK

Runtime, compiler, simulator

02

DEVELOPER SYSTEM

Cards and reference appliances

03

LFP EDGE

Modules, chiplets, licensable IP

04

LFP DATA CENTER

PCIe accelerators and scale-out systems

INITIAL WORKLOAD CANDIDATES

Time-series prediction • industrial anomaly detection • sensor fusion • adaptive control • cybersecurity • financial prediction • optimization

J4 wins where infrastructure is the bottleneck

The entry wedge is narrower than the total market—and large enough to support a semiconductor platform.

\$487B

2026 AI infrastructure

>\$1T

2029 AI infrastructure

\$22.7T

Enterprise AI

>\$40T

Humanoid robots

Win the workloads where power, cooling, latency, or autonomy drive the buying decision.

ILLUSTRATIVE YEAR 5 TARGET

\$650M

≈0.23% of the 2026 intelligent data-center silicon segment

Design partners create the shortest path to production

Start with workloads that have clear objective functions, high compute cost, and a measurable power or latency constraint.

01

DESIGN PARTNERS

3–5 anchor customers

02

SDK EVALUATION

Baseline vs. J4

03

FPGA / EMULATION

Kernel proof

04

DEV CARD

Customer pilots

05

PRODUCTION

Chip sales / licensing

INITIAL BEACHHEADS

Industrial automation • telecom / edge infrastructure • aerospace / defense • energy-constrained enterprise AI • colocation / data-center operators

Revenue expands as customers move from evaluation to deployment

Early revenue comes from paid integrations and development systems; long-term value comes from chips, software, and IP.

Design partnerships / NRE	Paid workload integration
Development systems	Cards, appliances, reference platforms
LFP chips + modules	Fabless semiconductor sales
IP + chiplet licensing	Upfront license + royalty
SDK + enterprise support	Subscriptions and support

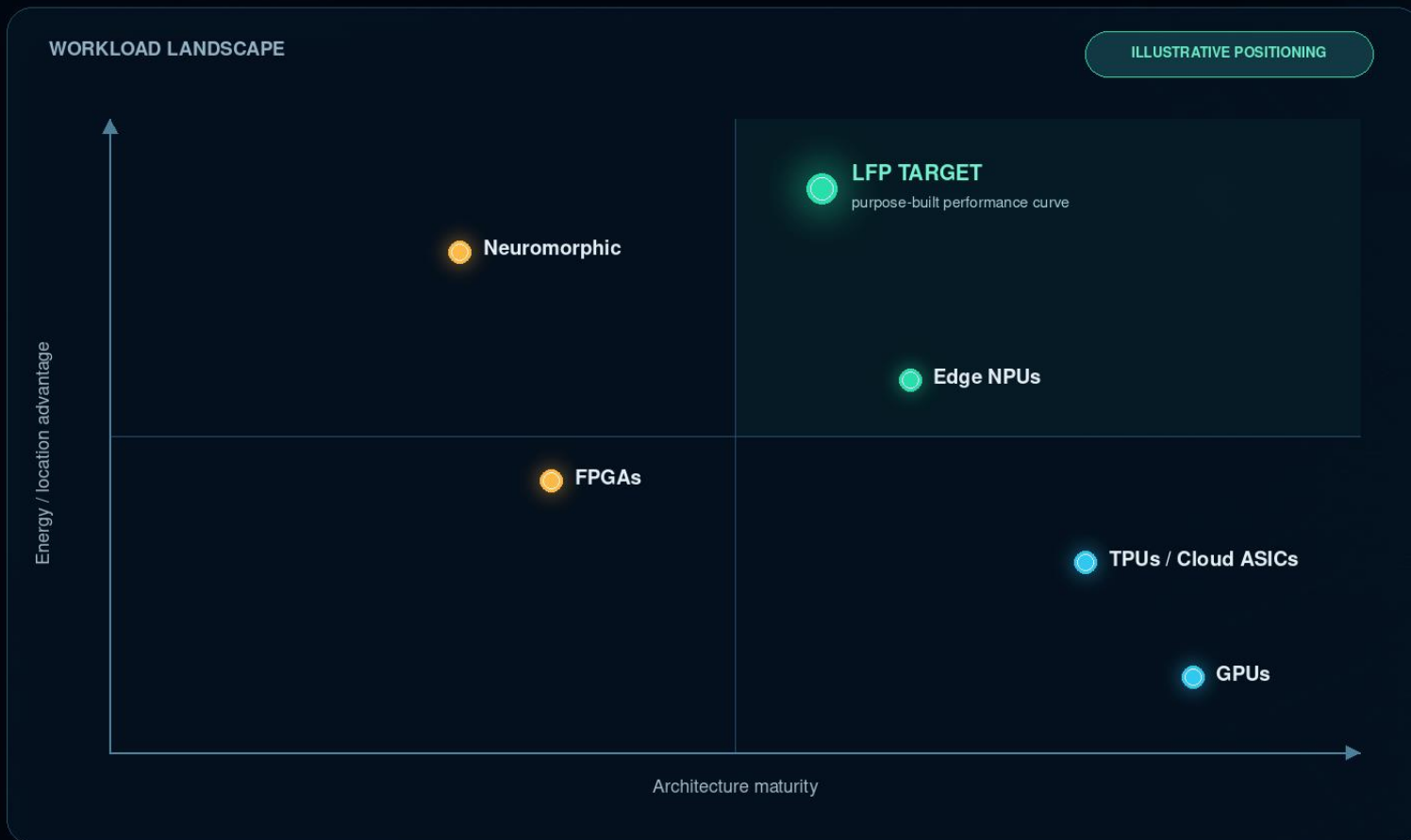
FABLESS

Outsourced
fabrication, packaging,
and test

Software + IP
grow margin over time

Win on workload economics

LFP does not need to beat every accelerator on every benchmark. It must dominate workloads where watts, cooling, autonomy, or latency drive the buying decision.



J4 DIFFERENTIATION

- Different learning model
- Purpose-built hardware
- Edge-to-data-center path
- Software + IP moat

Win where watts and milliseconds matter.

THE MOAT IS CUMULATIVE

IP and defensibility span method, silicon, and ecosystem

The learning method is the starting point. The implementation stack compounds the difficulty of replication.

PUBLISHED PATENT APPLICATION

WO2024258795A1
US20240419998A1

Legal status must be confirmed by patent counsel.

- 01 LFP instruction set + processing elements
- 02 SPDE assembly fabric + pruning hardware
- 03 Compiler, runtime, profiler, SDK
- 04 Benchmarks, workload corpus, customer co-design
- 05 Trade secrets + implementation know-how

A patent estate plus an implementation ecosystem that is difficult to replicate.

FOUNDATION COMPLETE; PROOF AHEAD

The next stage converts internal work into investor-grade evidence

The architecture is documented. The next proof points must be reproducible, independent, and customer-relevant.

TODAY

Published patent application

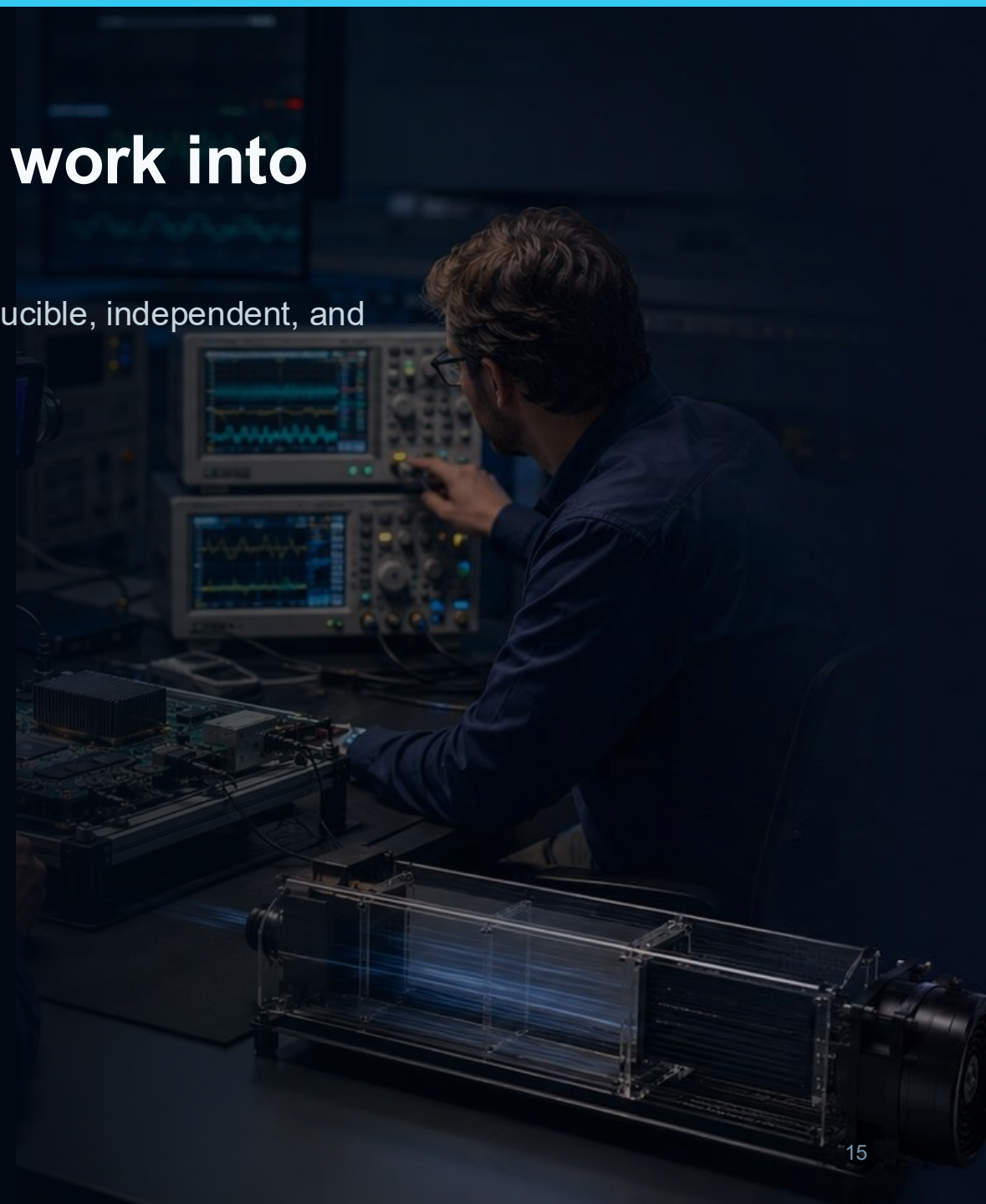
Software implementation

Energy reductions on selected computations

Defined LFP processor concept

NEXT PROOF POINTS

Reproducible benchmarks • independent energy and accuracy validation • FPGA kernels and scale-out model • 3–5 paid design partners • first LFP silicon



The financing assembles a semiconductor company around the inventor

The core method and technical vision are in place. The next hires convert them into commercial silicon and a platform business.

JEFF GLICKMAN

Founder + inventor

50 years

working on AI

Foundational ML patents

Core learning method

Technical vision + IP

SILICON LEADERSHIP

Build the machine

Chief silicon architect

RTL + verification

Physical design + packaging

Board / advisory priorities: advanced processor architecture • foundry / packaging • enterprise AI • edge systems • IP licensing • semiconductor finance

PLATFORM BUSINESS

Build the company

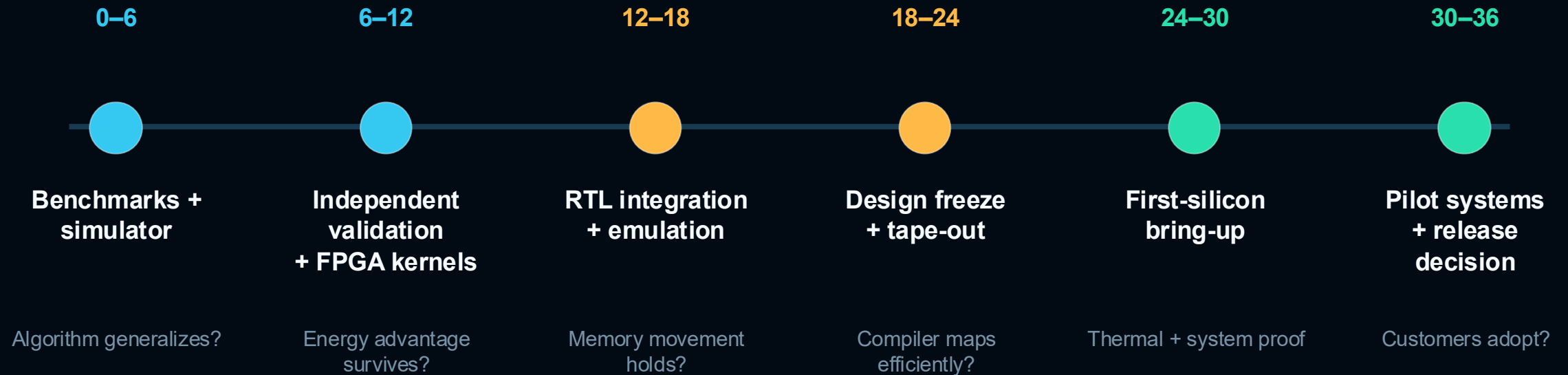
Compiler / runtime

BD + design partners

Finance, IP, operations

The 36-month plan sequences proof, integration, and first silicon

Capital deployment is gated by whether the algorithm generalizes, the energy advantage survives system effects, and customers adopt the model.



Each phase earns the next.

Revenue shifts from NRE to silicon, software, and licensing

A fabless model can expand gross margin as the product mix moves from evaluation to platform revenue.

REVENUE (\$M)



Management case (\$M)	Y1	Y2	Y3	Y4	Y5
Revenue	\$0	\$8	\$60	\$260	\$650
Gross margin	—	25%	48%	60%	64%
EBITDA	(\$39)	(\$46)	(\$35)	\$61	\$266

Illustrative management targets; not audited forecasts.

THE FINANCING REQUEST

\$150M to reach first-silicon proof and customer pilots

The structure is milestone-based: validate the method before committing the largest silicon dollars.

\$150M

RECOMMENDED TRANCHES

\$35M initial close

\$65M validation tranche

\$35M tape-out tranche

\$15M production reserve

Silicon architecture / RTL / verification	\$45M
EDA, IP, foundry, tape-out, packaging NRE	\$30M
Compiler, runtime, SDK, simulator	\$25M
FPGA, bring-up, development systems	\$15M
Design partners, pilots, market entry	\$12M
IP, security, compliance, legal	\$8M
Working capital + contingency	\$15M

Deliverable: a production-ready semiconductor company with first-silicon proof and customer pilots.

Generative AI cannot scale indefinitely by adding more power, more water, and more distance.

01

STOCHASTIC LEARNING

The method

02

LOGIC FIELD PROCESSOR

The machine

03

\$150M FINANCING

The platform build

Massive intelligence—wherever it is needed.

Legal disclosure

This communication is for informational purposes only and does not constitute an offer or solicitation to sell shares or securities in the Company or any related or associated company or any securities issued by funds managed by the Company. Any such offer or solicitation will be made pursuant to the applicable confidential Offering Memorandum and in accordance with the terms of all applicable securities and other laws. None of the information or analyses presented are intended to form the basis for any investment decision, and no specific recommendations are intended. Accordingly, this communication does not constitute investment advice or counsel or solicitation for investment in any security.

This communication does not constitute or form part of, and should not be construed as, any offer for sale or subscription of, or any invitation to offer to buy or subscribe for, any securities, nor should it or any part of it form the basis of, or be relied on in any connection with, any contract or commitment whatsoever. The Company expressly disclaims any and all responsibility for any direct or consequential loss or damage of any kind whatsoever arising directly or indirectly from: (i) reliance on any information contained in this communication, (ii) any error, omission or inaccuracy in any such information or (iii) any action resulting therefrom.

The data in this communication is unaudited. The statements contained herein are not intended to and do not constitute an opinion as to any tax or other matter. They are not intended or written to be used, and may not be relied upon, by you or any other person for the purpose of avoiding penalties that may be imposed under any Federal tax law or otherwise.

This material (including any attachments) contains information which is confidential and may also be privileged. This material is for the exclusive use of the intended recipient(s). If you are not the intended recipient(s), you are strictly prohibited from distributing, copying, disclosing or using this material (in whole or in part). If you have received this information in error, please notify us by email (info@j4.ai) or by telephone (425) 485-5600 and then destroy the material together with any copies of it.

All e-mail sent to or from J4.AI may be retained, monitored and/or reviewed. This communication does not constitute investment advice or a recommendation. J4.AI is a division of J4 Capital LLC.