



WE'RE HIRING · ENERGY INTELLIGENCE

FOUNDING ML PLATFORM ENGINEER

San Francisco · Hybrid · Full-time

Reports to: Lead, Energy Intelligence — Energy Software & Optimization

The new measure of energy.

Our data is rich rather than petabyte-scale. The hard problems here are messiness, freshness and trust — not raw volume. And every one of these inputs carries direct P&L and reliability consequences.

OWN
THE
SPINE.

THE ROLE

A founding, 0→1 build. You define what everyone else builds on.

You design and own the full data and ML spine the entire Energy Intelligence org runs on.

Joulent is a full-stack energy technology company built to deliver reliable, multi-gigawatt power at the speed and scale the AI era demands. Our modular, Across-the-Meter™ approach integrates generation, storage and advanced controls to power industrial hyperscalers. Our co-located power foundries bypass the constraints of the grid and bring electricity directly to the data center.

The Energy Intelligence team turns data into decisions: load- and market-facing forecasting across short- to long-term horizons, built on a shared, agent-native platform spine. This role runs from data engineering — ingestion, pipelines, feature store, quality and lineage — through the full MLOps lifecycle of training, serving, monitoring, and the agent and model gateway.

You move fast, ship to production, and treat AI agents as force-multipliers. If you want to build something real that helps power gigawatts of compute for the AI era, this is the seat.

BACKED BY

Engine No. 1
GE Vernova

SCALE

4 GW by end of 2027

Enough to power 3 million homes.

FORECAST HORIZONS

**Intra-hour,
day-ahead, seasonal**

WHAT YOU'LL DO

Three modes, one owner.

Builder Mode — stand up the spine

- Build production data pipelines and ML processes from scratch. You define the frameworks, SLAs and standards everyone else builds on.
- Own data ingestion end to end: market data (ISO/RTO feeds, LMPs, nodal price history), weather, interconnection-queue and plant/asset telemetry — managing quality, lineage and freshness.
- Build the feature store and simulation data plumbing, prevent training/serving skew, and ship AutoML pipelines so engineers train and deploy without rebuilding the plumbing each time.

Operator Mode — reliable, fast and cheap

- Own the ML lifecycle end to end: training orchestration, serving, registry, CI/CD, monitoring and drift detection.
- Stand up high-throughput, low-latency serving for intra-hour, day-ahead and seasonal forecasts, with metric-aware alerting.
- Define SLOs and dataset SLAs across freshness, quality and lineage; support make-vs-buy across data feeds and the MLOps stack.
- Balance performance, cost and reliability — and make the calls explicit and defensible.

AI-Native Mode — agents as force-multipliers

- Own the agent platform and model gateway — orchestration, routing, caching and batching — with token budgets, guardrails and observability across the model and agent fleet.
- Set the eval gates that decide which models and agents ship. Back-tests and evals are first-class here, with P&L and reliability on the line.

WHAT YOU'LL BRING

How you work matters as much as what you've shipped.

01

AI-native by default

You work alongside coding and analysis agents as a force-multiplier, not a novelty, and you stay current as the tooling moves. You raise the standard for how the team builds with AI.

02

High agency and urgency

You define requirements rather than wait for them. You know when the 60%-right-now answer beats the 95% answer in three days.

03

Cross-functional drive

You align engineers, data vendors and external partners without needing authority.

THE TECHNICAL BAR

What we need to see evidence of.

- 5+ years in a full-stack data and ML platform or production engineering role, having built production data pipelines and ML serving.
- Expert SQL and strong Python, with software-engineering rigor: CI/CD, containers and infrastructure-as-code.
- Hands-on experience running production data and ML pipelines end to end with orchestration tools like Airflow, across batch and streaming workloads on a modern cloud data platform.
- A track record of preventing training/serving skew, with data quality, lineage and observability built into your pipelines.
- Experience operating model serving, experiment tracking and drift monitoring in production.
- B.S. or M.S. in Computer Science, Machine Learning, Data/Software Engineering or a comparable field — or equivalent experience.

BONUS POINTS

Not required. Genuinely valued.

- Energy and market data: ISO/RTO, LMPs, nodal pricing, weather, interconnection-queue feeds. Time-series data at scale. Graph data and embeddings.
- LLM and agent ops: model gateways and routing, prompt and eval frameworks, token-cost optimization.
- Production cost modeling and simulation (PLEXOS, Aurora), and inference efficiency — quantization, distillation, transfer learning — for latency and cost budgets.

WHY YOU'LL LOVE IT HERE

A true inflection moment in the company's trajectory.

FOUNDATIONAL

We're moving deeper into the energy supply chain, and you get to build the data plumbing that makes it happen.

OWNERSHIP

A small, seasoned team where you own major parts of the system end to end, directly influencing the bottom line.

SCALE THAT MATTERS

A well-funded mission powering multi-gigawatt, national-scale infrastructure for AI-scale compute.



THE NEW MEASURE OF ENERGY.

HOW WE INTERVIEW

AI tools are allowed throughout our process, apart from one short fundamentals round — because that is how the job actually works. The take-home is capped at four hours and is paid.

APPLY

Sandy Suresh
Lead, Energy Intelligence
sandy.suresh@joulent.com