

Unit Economics of a Curve-Digitisation SaaS for Subsurface Teams

There is a well-worn failure mode in applied machine learning: a team builds a genuinely good model, digitises a hard corpus that nobody else can, and then discovers the economics of selling it were never worked out. This whitepaper is a product-and-go-to-market reading of a vertical curve-digitisation SaaS for subsurface teams, written from the position of having built the digitiser and then having to price it. It is deliberately not a re-telling of the model architecture, which we have published elsewhere; it is about the numbers a business runs on. The argument has one spine. A vertical SaaS becomes attractive only when the price a seat pays sits far above the marginal cost of serving that seat, and only when a disciplined market funnel keeps the revenue target honest rather than a number a slide deck asserts. We ground it in the sourced figures from one reference programme: a seat priced at 1,200 USD per user per year, a marginal serving cost derived from a GPU that rents for 750 to 1,800 EUR per month, a total market measured at 134B USD, a serviceable slice of 6.7B USD, a 3 percent obtainable share by end of year five reaching a 180M USD revenue projection against 6M USD in the first twelve months, a 3M USD seed with an 18-month runway, and two build tracks, an accelerated one at 180,000 EUR and a standard one at 100,000 EUR. Three instruments carry the argument: the contribution margin per seat, the seed as a fixed monthly burn ceiling, and the delivery-track fork where speed to revenue repays its own premium. The point is that the research being good is a precondition, not the business case, and the business case has to be built on numbers that survive division.

The EarthScan Team

Research, engineering, and field

VERTICAL-SAAS / UNIT-ECONOMICS / GROSS-MARGIN / MARGINAL-COST / GO-TO-MARKET / CURVE-DIGITISATION

CONTENTS

01	Why this document is about arithmetic, not architecture
02	The unit that repeats
03	What one seat actually costs to serve
04	The margin is necessary, not sufficient
05	The funnel that keeps the target honest
06	Three million over eighteen months is a spending rate
07	Why the ceiling and the margin talk to each other
08	Two ways to build the same product
09	When the standard track is still the right call
10	How the pieces constrain each other
11	What this looks like for a subsurface team specifically
12	Limitations
13	References
14	Get the full whitepaper

”

The model being good is a precondition, not the business case. A digitiser can hit the published accuracy and still lose money on every seat if the serving cost was never worked out. The business case is a small number of divisions, and each one has to survive contact with the real figures.

”

THE TWO QUESTIONS

Good research and a good business are not the same question

Why this document is about arithmetic, not architecture

We built a curve-digitisation model for subsurface teams and published its architecture and its accuracy in a peer-reviewed venue [1]. That document answers one question: does the model work. This document answers a different one that the first cannot touch: does the thing built around the model make money per seat, and is the plan to grow it honest. These are separate questions, and a team can pass the first and fail the second without noticing, because the accuracy figure feels like a business result and is not.

The reason to keep them apart is that they fail for unrelated reasons. A model fails on a benchmark. A business fails on a division: the price a customer pays turns out to be smaller than the cost of serving them, or the revenue target on the slide turns out to require a share of the market nobody has ever reached, or the money raised turns out to buy fewer months than the build takes. None of those failures shows up in an F1 score. All of them show up the moment you write the arithmetic down and refuse to round in your own favour.

So this is a product-and-go-to-market whitepaper, written from the position of having built the digitiser and then having to price it, staff it, and raise against it. It is grounded in the sourced figures from one reference programme, anonymised throughout: a digitisation effort for a Texas onshore operator, whose engagement archive supplies the seat price, the serving-compute rents, the market sizing, the seed terms, and the two build tracks we will read. Every number in the argument traces to that archive or to the published sources in the references, and where an input is a modelling assumption rather than a sourced figure we say so on the instrument and in the prose.

THE WHOLE BUSINESS CASE, AS SIX SOURCED NUMBERS

\$1,200

Seat price, per user per year

750-1800

GPU rent, EUR per month, the serving-cost source

\$134B / \$6.7B

Total market and serviceable slice

\$3M / 18mo

Seed ask and the runway it funds

The rest of the document takes those numbers and does the divisions in the open. First, the margin per seat, which is the question of whether the business works at all. Then the seed as a burn ceiling, which is the question of whether the build was sized to the money. Then the delivery-track fork, which is the question of whether it is worth paying to arrive sooner. The market funnel sits underneath all three, because it is what decides whether the seat count the divisions assume is a plan or a wish.

The unit that repeats

A subscription business is judged on the unit that repeats, not on the total it hopes to reach [3]. For a per-seat SaaS the unit is one seat: the annual revenue it returns and the direct cost of serving it for that year. If that unit is healthy, growth compounds it; if the unit loses money, growth multiplies the loss, which is the specific way a business with good-looking topline and bad unit economics dies. So the first thing to establish is not the market size or the revenue projection. It is whether one seat, served for one year, contributes more than it costs.

That is a favourable comparison for a curve-digitiser, for a structural reason worth stating before the numbers. The expensive part of a machine-learning product is building and training the model, and that cost is paid once. Serving is cheap and gets cheaper per unit as volume rises, because one trained model serves every log rather than being rebuilt for each one. The consequence is that the marginal cost of a curve-digitiser is dominated by serving compute, which is small, and the fixed cost is a one-off the seat price does not have to carry per unit. The margin question, then, is whether the seat price clears a serving cost measured in fractions of a cent per log. It does, and by a wide margin, which is the first instrument.

THE MARGIN

A seat pays 1,200 a year; serving it costs a fraction of that

What one seat actually costs to serve

The serving cost of a seat is built from the ground up out of sourced quantities. The model serves on a rented GPU. The engagement archive records two rentable tiers, 750 EUR per month for a high-end card and 1,800 EUR per month for an advanced one. That monthly rent, divided into an hourly rate, is the numerator of the per-log serving cost. The denominator is throughput: the number of logs the served model clears in an hour. Divide the hourly rent by the logs cleared that hour and you have the compute cost of digitising one more log, and it lands in fractions of a cent, because a rented GPU hour is a few tenths of a euro and a served model clears logs quickly.

The seat's yearly serving cost is that per-log cost times the number of logs one user runs in a year. This is where the structure of the product does the work. Because a single trained model serves every scan width, the per-log cost does not rise with the seat's usage, so the seat's yearly serving cost grows linearly in usage while the price it pays stays fixed at 1,200 USD. Even a heavy user, running thousands of logs a year, imposes a serving cost that is a small fraction of what the seat pays. The gap between the two is the gross contribution margin, and it is the whole business.

The three inputs we can vary honestly are the GPU tier, the seat's annual usage, and, buried in the derivation, the throughput and the rented-month-to-hour conversion. The tier and usage are levers on the instrument; the throughput, the hours-per-rented-month divisor, and the usage band are modelling inputs flagged as illustrative, because the archive gives us the tier prices and the seat price but not a measured serving throughput at production scale. What the instrument argues does not depend on the illustrative inputs being exact. It depends on the shape: a fixed price column that a growing cost column cannot come close to reaching at any realistic usage.



The gross-margin argument for a vertical curve-digitisation SaaS, stated as the one comparison the business rests on: the 1,200 USD a seat pays each year against the marginal compute it costs to serve that seat. Lever A toggles the rented GPU tier, 750 EUR per month for a high-end card or 1,800 EUR for an advanced one, which sets the hourly rent and therefore the serving cost of one more digitised log. Lever B drags the logs one seat runs per year, the term that scales the seat's serving cost; because one trained model serves every scan, the per-log cost stays flat, so the cost column grows linearly while the price column stays fixed at 1,200. The orange bracket is the only element that argues: the contribution margin, the gap between price and serving cost that has to stay wide for the business to work, and at any realistic serving throughput it does. The seat price and the two GPU tier prices are sourced; the hours-per-rented-month divisor, the logs-per-hour serving rate, and the logs-per-seat usage band are illustrative serving inputs, and this is an engineering unit economics view, not an accounting statement.

The ledger reads the margin directly. The teal column is the 1,200 USD a seat pays, fixed. The dim column is the serving cost that seat imposes, growing as you drag its yearly usage up. The orange bracket between them is the only element that argues, and it is the contribution margin: at any realistic serving throughput it stays wide, because the serving cost of even a heavy seat is a small fraction of the price. Toggle to the advanced GPU tier and the cost column rises a little, because the rent went up, and the margin barely moves, because the rent was small relative to the price to begin with. That is the thesis of the whole whitepaper stated as one picture: a vertical curve-digitisation SaaS is attractive precisely because the price sits far above the marginal serving cost per log, and it stays attractive as usage grows because one model serves everything.

THE POSTURE ON MARGIN

Price the seat against the marginal serving cost, not against the cost of building the model. The build is a one-off that the seat price does not carry per unit; the serving cost is what recurs, and for a curve-digitiser it is small and roughly flat in usage. If a vertical SaaS cannot show a wide gap between seat price and marginal serving cost, no amount of topline growth will fix it, because growth multiplies the unit.

The margin is necessary, not sufficient

A wide per-seat margin is the entry ticket, not the whole game. It says the unit is healthy, which means growth will compound rather than dilute it. It says nothing about whether you can acquire the seats, whether the money you raised sizes the build, or whether the revenue projection you are selling against is reachable. Those are the next three questions, and the first of them is the one a wide margin can seduce a team into skipping: how big is the reachable market, really, and is the revenue target a share of it that anyone has ever obtained.

The funnel that keeps the target honest

The reference deck sizes the market in the conventional three steps. The total addressable market for the transactions the product touches is 134B USD. The serviceable slice, the oil-and-gas technology market the product can actually address, is 6.7B USD, which is about 5 percent of the total. The obtainable share the plan commits to is 3 percent of that serviceable slice by end of year five, which is the 180M USD revenue projection, against 6M USD in the first twelve months after the seed. We have written a dedicated instrument for that walk-down elsewhere and will not repeat it here; the point that matters for the margin argument is narrower.

It is this. The 180M USD figure is not a number to lead with; it is a number to check against the seat price. At 1,200 USD per seat, 180M USD of revenue is 150,000 seats, and 6M USD in year one is 5,000 seats. Those are the real targets, expressed in the unit the business actually sells. A revenue projection stated in dollars can float free of reality; the same projection stated in seats has to answer for whether a sales organisation can close 5,000

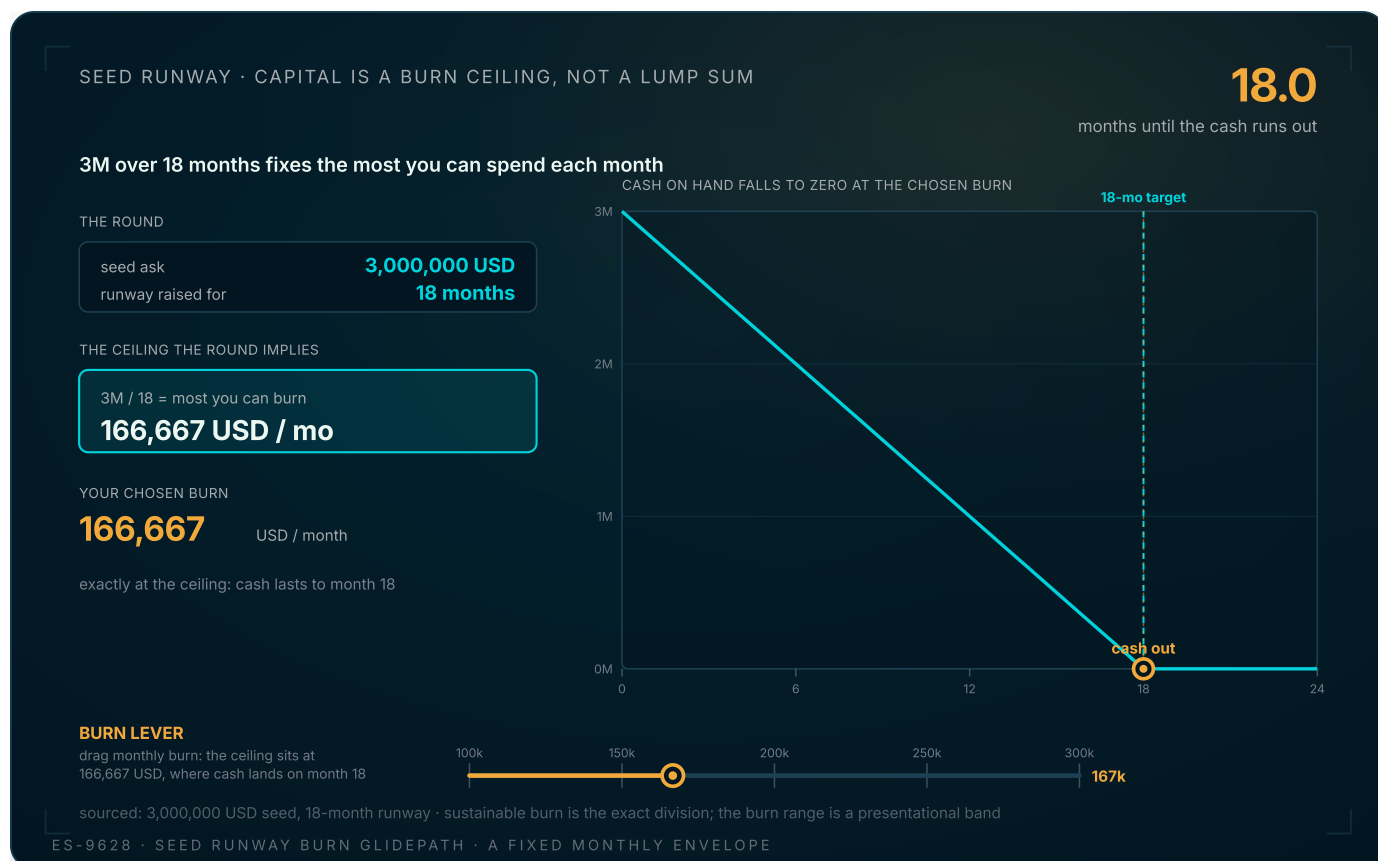
seats in a year and grow that to 150,000 by year five. Keeping the target in seats, and demanding that the funnel from market size down to obtainable share be a share someone has plausibly reached, is what stops the projection from becoming a spreadsheet fantasy. The wide margin makes each seat worth selling; the honest funnel makes the seat count a plan rather than a wish.

A seed is a burn ceiling, not a lump sum

Three million over eighteen months is a spending rate

The seed round is 3M USD, raised for an 18-month runway. The mistake a team makes with a number like that is to read it as a pile of money to spend as opportunities arise. It is not. A seed raised for a fixed runway is a fixed maximum monthly burn, because the capital divided by the runway is the most you can spend each month and still reach the end of the runway with the company alive. Three million over eighteen months is 166,667 USD per month, and that figure, not the three million, is the constraint every hiring and build decision has to fit inside.

This reframing matters because it converts a vague sense of having money into a hard ceiling that the build plan must respect. If the plan implies a burn above the ceiling, the runway ends before the milestones the round was raised to reach, and the company is raising its next round from a position of weakness or not raising it at all. If the plan sits below the ceiling, there is slack to absorb the things that always go wrong. The seed does not fund a wish list; it funds a rate, and the rate is what the second instrument makes visible.



The seed round read as a hard monthly envelope rather than a lump sum. A 3,000,000 USD seed raised for 18 months of runway sets a fixed maximum sustainable burn of 166,667 USD per month, the exact division of the two. Drag the burn lever and the glidepath, cash on hand falling in a straight line as burn times months eats the capital, shows where the money runs out. Below the ceiling the line reaches the 18-month target with cash to spare; at the ceiling it lands exactly on month 18; above it, the orange runway-end marker slides left of the target and the round fails to reach the milestone it was raised to fund. The orange marker is the only element that argues: the month the cash runs out, plotted against the 18-month target. The seed size and the runway are sourced from the deck; the sustainable burn is their exact quotient, and the burn-lever range is a presentational band around it.

The glidepath draws the capital falling to zero as burn times months eats it. The teal line is the cash on hand; the teal dashed line is the 18-month target the round was raised to reach; the orange marker is the only element that argues, and it is the month the cash runs out. Drag the burn below the 166,667 USD ceiling and the line reaches month 18 with cash to spare. Drag it to the ceiling exactly and the line lands on month 18, which is what the round was sized for. Drag it above and the orange marker slides left of the target, and you can watch the round fail to buy the runway it promised. The instrument argues one thing: the seed sets a burn ceiling, and the ceiling is what keeps the revenue plan honest, because a plan that requires spending above the ceiling is a plan that runs out of money before it can be tested.

Why the ceiling and the margin talk to each other

The burn ceiling and the per-seat margin are not separate facts; they constrain each other. The wide margin means every seat sold reduces the net burn, because the revenue a seat returns is almost all contribution. So the burn ceiling is not a fixed sentence; it softens as seats come on. But that only helps if the seats arrive inside the runway, which is why the timing of first revenue is the next question, and why paying to ship sooner is a real lever on survival rather than a vanity. A team that ships late burns through the runway with no offsetting revenue; a team that ships early starts converting the wide margin into cash while there is still runway to protect. That is the bridge from the money question to the delivery question.

"The seed did not buy us three million dollars of freedom. It bought us eighteen months at a hundred and sixty-seven thousand a month, and the only way to loosen that was to start selling seats before the runway ran down."

— FROM OUR OWN PLANNING NOTES



THE TIMING

Paying to ship sooner buys weeks of revenue

Two ways to build the same product

The engagement archive records two build tracks for the product. The accelerated track costs 180,000 EUR, staffs 6 FTE, and finishes in 16 weeks. The standard track costs 100,000 EUR, staffs 4 FTE, and finishes in 32 weeks. The accelerated track costs 80,000 EUR more and ships 16 weeks sooner. The naive reading is that the standard track is cheaper and therefore the safer choice for a company watching a burn ceiling. The correct reading is that the two tracks are not being compared on cost; they are being compared on when the revenue starts, and 16 weeks of earlier revenue at 1,200 USD per seat can be worth far more than the 80,000 EUR premium.

The reasoning is the payback question. Both tracks build the same product, so both earn the same revenue per seat once they ship. The only difference is the week the earning begins. The accelerated track starts earning 16 weeks before the standard one, and during those 16 weeks it accumulates revenue the standard track earns nothing against. The premium is repaid at the moment that head-start revenue clears 80,000 EUR, and whether that moment arrives inside the first year depends only on how many seats are live. Enough live seats and the premium repays itself well before the standard track has even shipped.

Paying 80k more to ship 16 weeks sooner buys 16 weeks of revenue

TWO BUILD TRACKS

Accelerated

180,000

EUR build cost

6 FTE

16 weeks

Standard

100,000

EUR build cost

4 FTE

32 weeks

premium paid

80,000 EUR

weeks shipped sooner

16 weeks

revenue per week

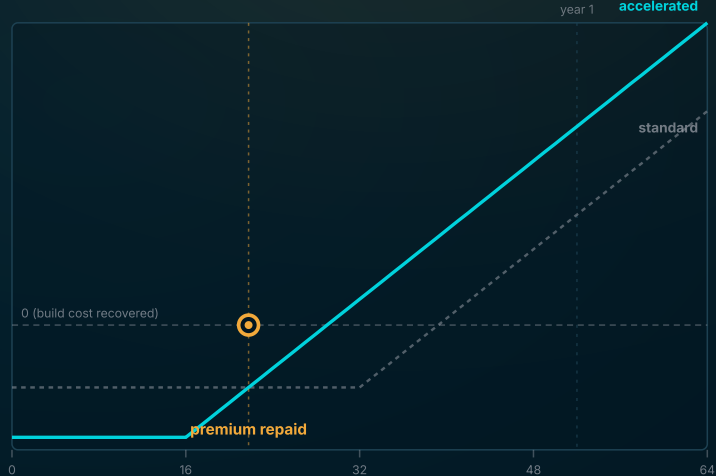
13,846 USD

at 600 live seats

premium repaid in year one

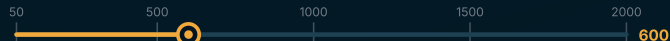
the head start pays for itself when seats are live

CUMULATIVE NET CASH · EACH TRACK EARN FROM ITS OWN FINISH



SEATS LEVER

drag live seats: both tracks earn the same per seat; only the start week differs



sourced: 180k EUR/6 FTE/16 wk vs 100k EUR/4 FTE/32 wk, 1,200 USD/seat/yr · the seats-live count is a presentational input

ES-8758 · DELIVERY TRACK PAYBACK FORK · SPEED TO REVENUE REPAYS THE PREMIUM

The build-track decision drawn as a time-to-revenue fork. The product can be built two ways: an accelerated track at 180,000 EUR with 6 FTE that finishes in 16 weeks, or a standard track at 100,000 EUR with 4 FTE that finishes in 32 weeks. The accelerated track costs 80,000 EUR more but ships 16 weeks sooner, and every earlier week is a week of seat revenue at 1,200 USD per user per year. Drag the live-seats lever and the two cumulative-net-cash curves accrue from each track's own finish week across a shared timeline; the orange marker sits where the accelerated track's head-start revenue has repaid its 80,000 EUR premium. The argument is single: paying the premium is rational exactly when the revenue the head start unlocks exceeds it, and enough live seats make that true well inside the first year. The track costs, FTE counts, timelines, and seat price are sourced from the engagement archive; the live-seats count is a presentational input, and this is an engineering timing view, not an accounting statement.

The fork plots each track's cumulative net cash from its own finish week. The teal line is the accelerated track, starting at minus 180,000 EUR and climbing from week 16. The dim dashed line is the standard track, starting at minus 100,000 EUR and climbing only from week 32. The orange marker is the only element that argues, and it sits at the week the accelerated track's head-start revenue has repaid its 80,000 EUR premium. Drag the live-seats lever up and the marker slides left, because more seats mean the 16-week head start is worth more per week, and the premium repays faster. At a serious seat count the premium is repaid inside the first year, often long before the standard track has shipped at all, which is the argument: paying to arrive sooner is rational exactly when the revenue the head start unlocks exceeds the premium, and a wide-margin product with live demand makes that easy to clear.

When the standard track is still the right call

The fork does not say always pay the premium. It says the premium is worth paying when the head-start revenue clears it inside a horizon you care about, and that condition can fail. If demand is thin at launch, few seats are live, the head start is worth little per week, and the premium takes a long time to repay, which argues for the cheaper track and its lighter draw on the burn ceiling. If the product is not yet validated and the risk is building the wrong thing fast, the standard track buys time to learn at a lower cost. The instrument is a decision aid, not a verdict: it tells you, at a given seat count, whether the premium repays inside the year, and it lets you see how many live seats it takes to make speed the right call. The honest answer depends on the demand you actually have, which loops back to the funnel: speed only pays when there are seats to earn against.

THE SYNTHESIS

The four questions are one system

How the pieces constrain each other

Read separately, the four questions look like a checklist. Read together, they are a single system where each answer sets the terms for the next. The per-seat margin decides whether the unit is worth selling at all; because it is wide, every seat sold is almost pure contribution. That wide margin is what makes the seed's burn ceiling survivable, because seats coming on soften the burn faster than they would for a thin-margin product. The burn ceiling is what makes the timing of first revenue matter, because a runway with no offsetting revenue is a runway spent; and the timing is what makes the delivery-track premium rational, because shipping 16 weeks sooner converts the wide margin into cash while there is still runway to protect. And the market funnel sits underneath all of it, because every one of these arguments assumes seats can be acquired, and the funnel is what turns the revenue target from a dollar figure into a defensible seat count.

The failure of any one piece propagates. A thin margin makes the burn ceiling unsurvivable no matter how the build is timed. A dishonest funnel makes the seat count a fantasy, and then the payback fork is repaying a premium against revenue that never arrives. A build sized above the burn ceiling runs out of money before the margin can prove itself. The reason to write all four down as arithmetic, and to refuse to round any of them in your own favour, is that they are load-bearing on each other, and a soft number in one place quietly weakens the argument everywhere.

What this looks like for a subsurface team specifically

Nothing in the argument so far is specific to subsurface work, which is deliberate: the unit-economics reasoning is the same for any vertical SaaS. What is specific is why the margin is so wide in this case, and it comes from the domain. A subsurface team sits on a corpus

of raster logs that is expensive to interpret by hand and that accelerates enormously once digitised, with the acceleration over manual interpretation large enough to be worth paying for on its own terms [2]. That means the value a seat receives is high, which supports the 1,200 USD price, while the cost of delivering it is dominated by cheap serving compute, which keeps the marginal cost tiny. The wide margin is not an accident of pricing; it is the domain's expensive manual baseline meeting the product's cheap automated serving, and the gap between them is the business.

The corollary is a warning. The same reasoning that makes the margin wide makes it tempting to skip the harder questions, because a wide margin feels like the business is already won. It is not. A wide margin with an unreachable seat target, or a build that overruns the burn ceiling, or a launch timed so late the runway is gone before revenue starts, still fails. The margin is the reason to do the work, not a substitute for it.

WHAT TO CARRY OUT OF THIS

- 1** The model being good is a precondition, not the business case. A curve-digitiser can hit its published accuracy and still lose money per seat if the serving cost was never priced. The business case is a small number of divisions, and each has to survive contact with the real figures.
- 2** A vertical SaaS works only when the seat price sits far above the marginal serving cost. At **1,200** USD per seat against a serving cost of fractions of a cent per log, the contribution margin is wide, and it stays wide as usage grows because one trained model serves every scan.
- 3** The revenue projection has to be read in seats, not dollars. **180M** USD by year five is **150,000** seats and **6M** USD in year one is **5,000** seats; a funnel from a **134B** USD market through a **6.7B** USD serviceable slice to a **3** percent obtainable share only holds if that share is one someone has reached.
- 4** A **3M** USD seed raised for **18** months is a burn ceiling of **166,667** USD per month, not a lump sum. Every build and hiring decision has to fit inside the rate, or the runway ends before the milestones the round funded.
- 5** Paying **80,000** EUR more to ship **16** weeks sooner buys 16 weeks of seat revenue. The premium repays itself when the head-start revenue clears it, which enough live seats make happen inside the first year, so speed to revenue is a survival lever, not a vanity.

Limitations

The per-seat serving cost in the margin ledger is built from sourced GPU tier prices, 750 and 1,800 EUR per month, and the sourced seat price of 1,200 USD, but the serving throughput in logs per hour, the hours-per-rented-month divisor that converts monthly rent to an hourly rate, and the logs-per-seat-per-year usage band are modelling inputs flagged as illustrative on the instrument, because the engagement archive supplies the rents and the price but not a measured production serving throughput. The argument the ledger makes, that the margin stays wide at any realistic usage, does not depend on those inputs being exact; it depends on the serving cost being small relative to the price, which the sourced rents and price already establish. The margin comparison places a EUR serving

cost against a USD price near currency parity for readability, which is a simplification, not an accounting treatment; a precise statement would carry the exchange rate and the fixed company costs the gross margin sits above. The seed glidepath treats burn as constant across the runway, which real burn is not; the sustainable monthly figure is the exact division of the sourced 3M USD seed by the sourced 18-month runway, and the burn-lever range around it is presentational. The delivery-track fork uses the sourced track costs, FTE counts, and timelines and the sourced seat price, but the live-seats count is a presentational input, and the weekly seat revenue is the annual price divided evenly across 52 weeks, which ignores ramp and churn; the fork argues the direction of the timing trade, not a precise payback date. The market funnel figures are the reference deck's own numbers; whether a 3 percent obtainable share is reachable is a judgement the arithmetic cannot settle, and the whitepaper's claim is only that the target should be checked in seats, not that the share is assured. Finally, the wide margin is specific to a domain where the manual baseline is expensive and the automated serving is cheap; a vertical where that gap is narrower would not support the same reasoning, and the method here is the reasoning, not the particular numbers.

References

- 1 Maiti, T., Patwardhan, N., Tambe, S. (2023). *VeerNet: a deep learning model for the automated digitization and interpretation of well logs*. MDPI Journal of Imaging, 9(7), 136. The published architecture whose serving compute sets the marginal cost this whitepaper prices against. <https://www.mdpi.com/2313-433X/9/7/136>
- 2 Koroteev, D., Tekic, Z. (2021). *Artificial intelligence in oil and gas upstream: Trends, challenges, and scenarios for the future*. Energy and AI. The context for why a subsurface team has a corpus worth digitising, and why the acceleration over manual interpretation is large enough to sell against. <https://www.sciencedirect.com/science/article/pii/S2666546820300033>
- 3 Skok, D. (2013). *SaaS Metrics 2.0: A Guide to Measuring and Improving What Matters*. For Entrepreneurs. The reference framing for reading a subscription business through the unit that repeats: the cost to acquire and serve a customer against the revenue that customer returns. <https://www.forentrepreneurs.com/saas-metrics-2/>

Get the full whitepaper

This page is the long-form summary. The complete whitepaper adds the per-seat margin table across GPU tiers and usage tiers, the seed glidepath under variable rather than constant burn, the full market funnel from 134B USD down to the 3 percent obtainable share with the seat-count translation at each step, and the delivery-track payback worked at several demand scenarios.

Unit Economics of a Curve-Digitisation SaaS for Subsurface Teams

Published August 2024. Generated 2026-07-22.

<https://www.earthscan.io/whitepapers/unit-economics-curve-digitisation-saas-subsurface-teams>

Copyright EarthScan. All rights reserved.