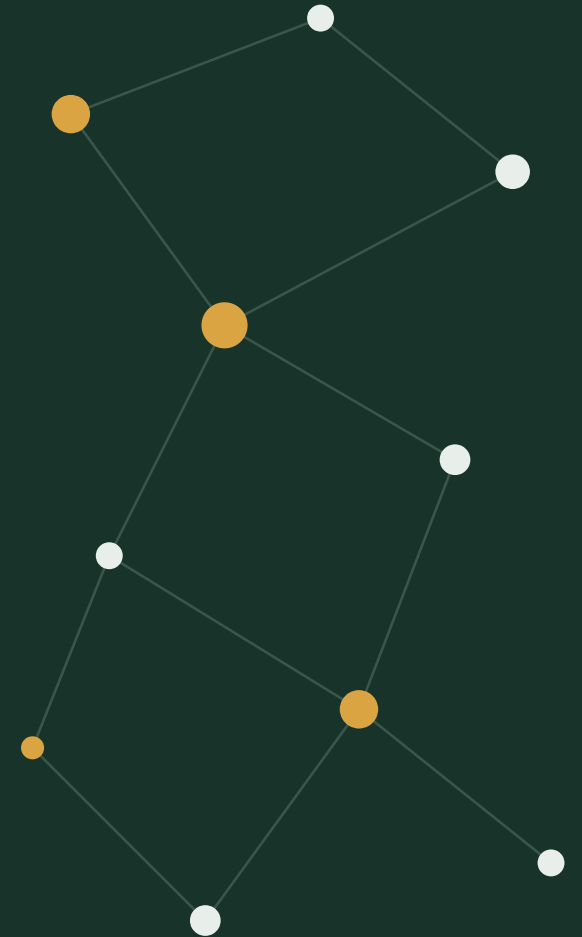


ITALIAN PORTFOLIO · STRATEGY · AUGUST 2026

One Graph, Three Doors

The durable asset under Paroletta, Filotto, and Salotto is the learner's known-word graph. Build the asset now. Monetize it later.

Dave Paola · paroletta.com · filottolang.com · salottolang.com



SITUATION

The science is settled, the market is enormous, and we are already live



Acquisition happens at $i+1$

Krashen's input hypothesis: you acquire language from input one step past what you personally hold. Not "A2 material," which is a population average. The whole field agrees on this. Almost nobody ships it.



The market pays for the average version

Duolingo: roughly \$750M revenue in 2024 and about 40M daily actives, built on fixed content tracks and a streak. The entire category personalizes to a CEFR band at best.



We are not pitching an idea

Three products live in production, a WhatsApp learner community, and a nightly generated-input loop proven on one learner for months. Everything on the next slides is running code and real telemetry.

LIVE TODAY

paroletta.com

daily word game: web, App Store, Android

filottolang.com

cloze drill live, 18,952-sentence corpus

salottolang.com

live bilingual captions, two versions deployed

WhatsApp group

the A1 to B2 learner community we launched into

Nightly email loop

generated chapters at measured coverage, $n=1$

Everyone sells content. Nobody measures the learner.

Personalized input requires knowing which words this learner already holds. Acquiring that knowledge is a cold-start bind with three walls:

1

Measure properly, lose the user

A real vocabulary assessment is hundreds of items. Nobody sits through that before seeing value.

2

Skip it, and you cannot personalize

Which is where the whole category lands. Hence CEFR bands: a blunt instrument standing in for a measurement nobody takes.

3

Borrow it, depend on a competitor

The only parties holding a good graph are the ones you would compete with. Clozemaster exposes no export API at all.

THE MOAT EVERYONE HAD IS GONE

LLMs regenerate any course, any story, any drill in an afternoon. Content is no longer defensible for anyone, including the incumbents who spent a decade building it.

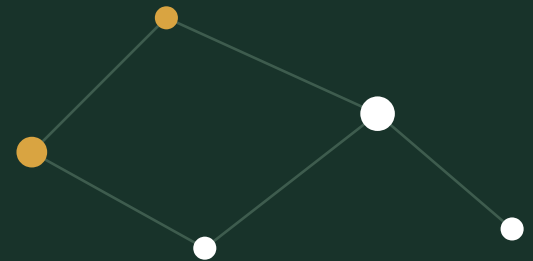
THE LEARNER PAYING THE PRICE

The Stranded A2: 600 to 1,500 words from Duolingo or Clozemaster and nothing to read with them. The category teaches the beginning of the language, forever, because the beginning is the only part that works without measurement.

THE QUESTION

How do you build something durable in language learning when content is free and the measurement everyone needs is locked up?

(And the sub-question every audience asks first: why won't Duolingo just do this? Slide 11.)



Make the measurement the activity. The graph is the asset.



1 • Every product is an instrument

Stop treating assessment as a tax paid before value. Ship things learners genuinely want to do, designed so every interaction emits vocabulary signal as a byproduct. The user experiences a game. We operate an instrument.



2 • The signal compounds into one graph

Each learner's known-word graph: which words they can produce, which they merely recognize, which they reach for and miss. It compounds daily, it cannot be copied, and it powers every surface at once.



3 • Build the asset, monetize later

The explicit goal. Revenue options multiply as the graph grows: the reader subscription, the graph as an API, adjacent languages. Choosing one today would shrink the asset to fit the business. Wrong order.

Four arguments follow, each with its evidence: the graph is a moat, the cold start dissolves, the portfolio feeds one graph, and speed is structural.

The known-word graph is a moat, not a feature



It compounds

Every session adds signal. A competitor launching today starts at zero against a learner's years of history with us.



It cannot be copied

Content can be regenerated by any competent model in an afternoon. A per-learner measurement history cannot be regenerated at all.



It transfers across surfaces

One graph powers reading, listening, drilling, and conversation. Each new surface both spends the asset and feeds it.

EVIDENCE

The incumbents already treat it as the crown jewels: Clozemaster holds exactly this graph and exposes no export API; LingQ and Anki come closest and stop at counts and cards. **Our 2026 competitive scan of generated-reader products found generation fully commoditized: measurement was the only defensible piece left on the board.**

The cold start dissolves when the game is the test

Day one needs no user data: sentences arrive in frequency order, the same first guess Clozemaster makes. By week two the learner's own play has replaced the prior. From there, every surface is a probe:

Surface	The learner experiences	The graph receives	Measured yield
Paroletta	A 30-second daily word game	Every typed guess: words produced unprompted, plus la domanda, 15 to 30 written words a day	About 100 distinct produced words per player per month
Filotto	A cloze drill and a nightly chapter written inside their vocabulary	Answered words, stored separately from ride-along words, from the first commit	18,952-sentence corpus live; the 15-minute daily habit
Salotto	A live conversation room with bilingual captions	Full transcripts, including the code-switch stall: the exact word they reached for and missed	200 to 400 distinct lemmas per person per session

A wrong guess carries more information than a right one: guess CARTA against answer BOCCA and we learn the player produces CARTA, was reaching for a five-letter noun, and does not hold BOCCA.

One graph, three doors today, each sized to its own cadence



Paroletta

PLAY • 30 SECONDS DAILY

Top of funnel. Anonymous, one tap, shareable with strangers. Emits produced guesses and daily written answers.



Filotto

TRAIN • 15 MINUTES DAILY

The habit. Identified, stateful, accumulating. Emits cloze answers; spends the graph on nightly generated chapters.



Salotto

SPEAK • SCHEDULED, SOCIAL

The richest sample. Live rooms for the community. Emits full conversation transcripts and reaching moments.



THE KNOWN-WORD GRAPH one shared measurement per learner, fed by every door, owned by us



What the graph buys: a chapter every night the learner can actually read, verified against their vocabulary before it sends • audio of text already read, for the walk • drilling at the edge of what they hold. Delivered by the rails that never miss: email, push, and the WhatsApp community. And the doors multiply: the candidate pipeline is on slide

15.

AI-native speed is structural, not heroic



3

products shipped to production in 18 days

1

builder, working nights and weekends

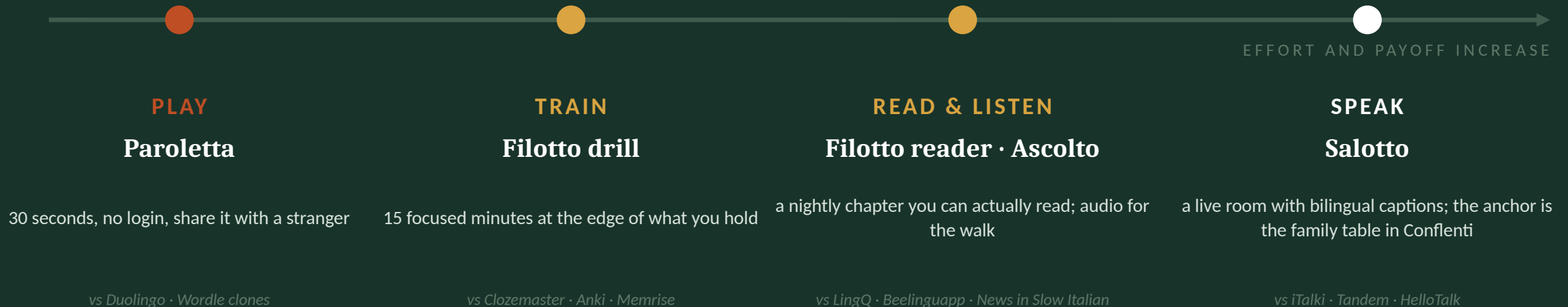
24h

from anomaly in telemetry to coordinate-level root cause and a designed fix

Call it 100x if you like. The precise version: **our experiment cadence is measured in days where the incumbents' is measured in quarters, and cadence compounds.** Duolingo is a gamification company that uses AI. This is an AI-native operation where one person directs a fleet of agents, and the org chart is a prompt.

The spectrum from playing to speaking

Learning a language is not one activity. It runs from play, which costs nothing to start, to speech, which costs the most and teaches the most. Incumbents park the learner at a single point. We built the whole spectrum on one graph.



Every incumbent owns one point on this line; none carries a measurement across them. Duolingo's bet: park the learner at PLAY forever. Ours: **a learner who moves along the spectrum stays for years, and every step makes the graph, and the product, better.**

“Why doesn't Duolingo just do this?”

Different learner

They monetize the beginner, forever. Our learner is the one they graduate and abandon: the Stranded A2 with 600 to 1,500 words and nothing to read. Serving them requires admitting the streak did not produce fluency.

Different asset

Their compounding asset is the streak and the engagement machine around it. Ours is the learner's lexicon. Telling a 900-day-streak user "we now measure what you actually know" is an awkward conversation for them and our entire premise.

Different economics

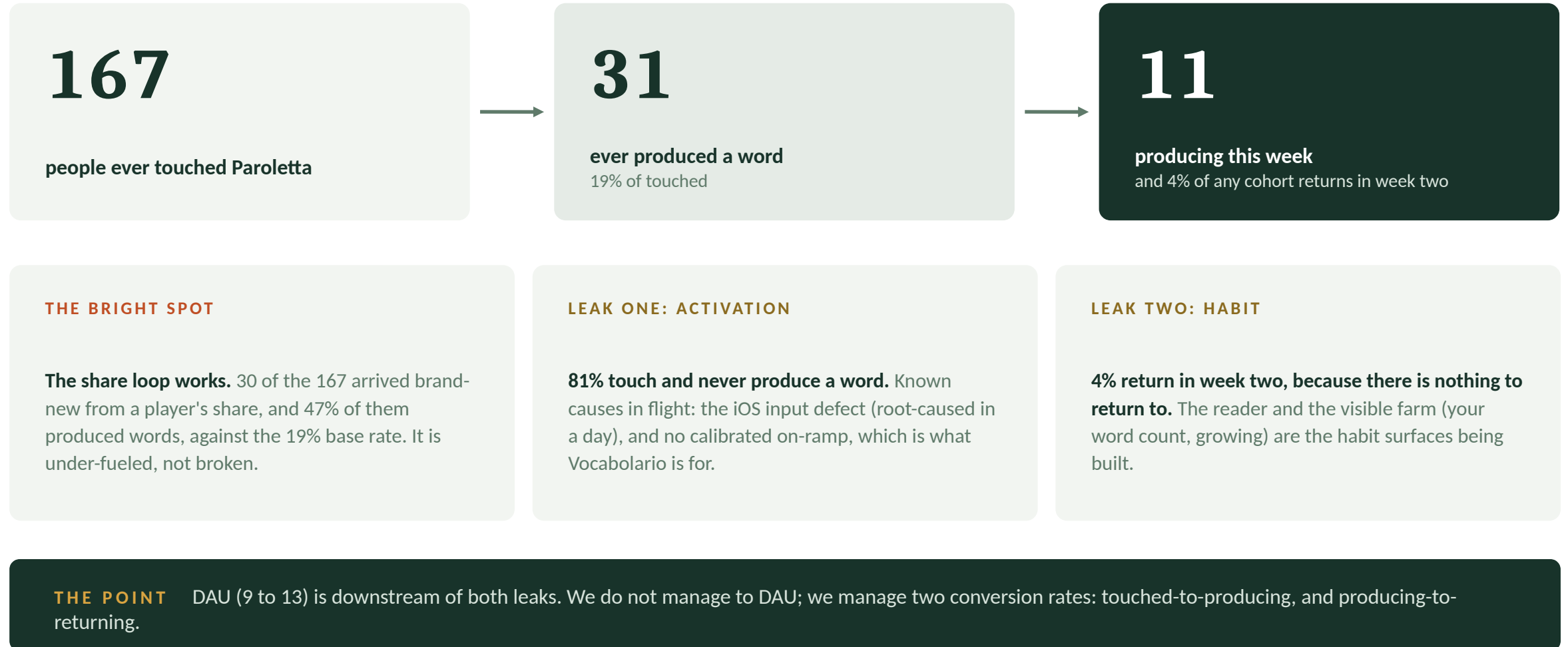
They carry the cost structure of a content company. We carry none: generation is free by assumption, so 100% of effort goes to measurement and delivery. Their moat evaporating is our starting condition.

Different speed

They are a public company that A/B tests features over quarters. We ship whole products in days and run the portfolio as a set of falsifiable experiments with predictions written before the data arrives.

The honest caveat: they have the data to build a graph tomorrow. Our bet is that they will not want to, for the three reasons above, and that by the time they do, per-learner history is the one thing they cannot backfill.

The funnel is measured, and it leaks in two named places



The scoreboard: one north star, a KPI per door

NORTH STAR distinct produced words added to learner graphs per week. Everything else is diagnostic.

Door	KPI to move	Baseline, Aug 20
Paroletta	Produced words per player per month; week-2 return; touched-to-producing rate	~100 words; week-2 return 4%; 19% of touched ever produce
Filotto	Drill answers per learner per week; nightly chapter open rate	Drill and corpus live; reader pre-launch, no baseline yet
Salotto	Distinct lemmas per person per session; sessions held per month	200 to 400 lemmas measured; 0 sessions currently scheduled
Distribution	Share arrivals per week and their producing conversion, by artifact and channel	~2 arrivals per week, converting at 47%, twice the base rate

Live today: the "Graph feed rate" trend in PostHog (weekly typed guesses, domanda answers, producing learners). **Filed:** a nightly graph_snapshot event from our own store, because the words themselves never reach PostHog; once it ships, true graph size becomes a chart instead of an inference.

The distribution plan, channel by channel

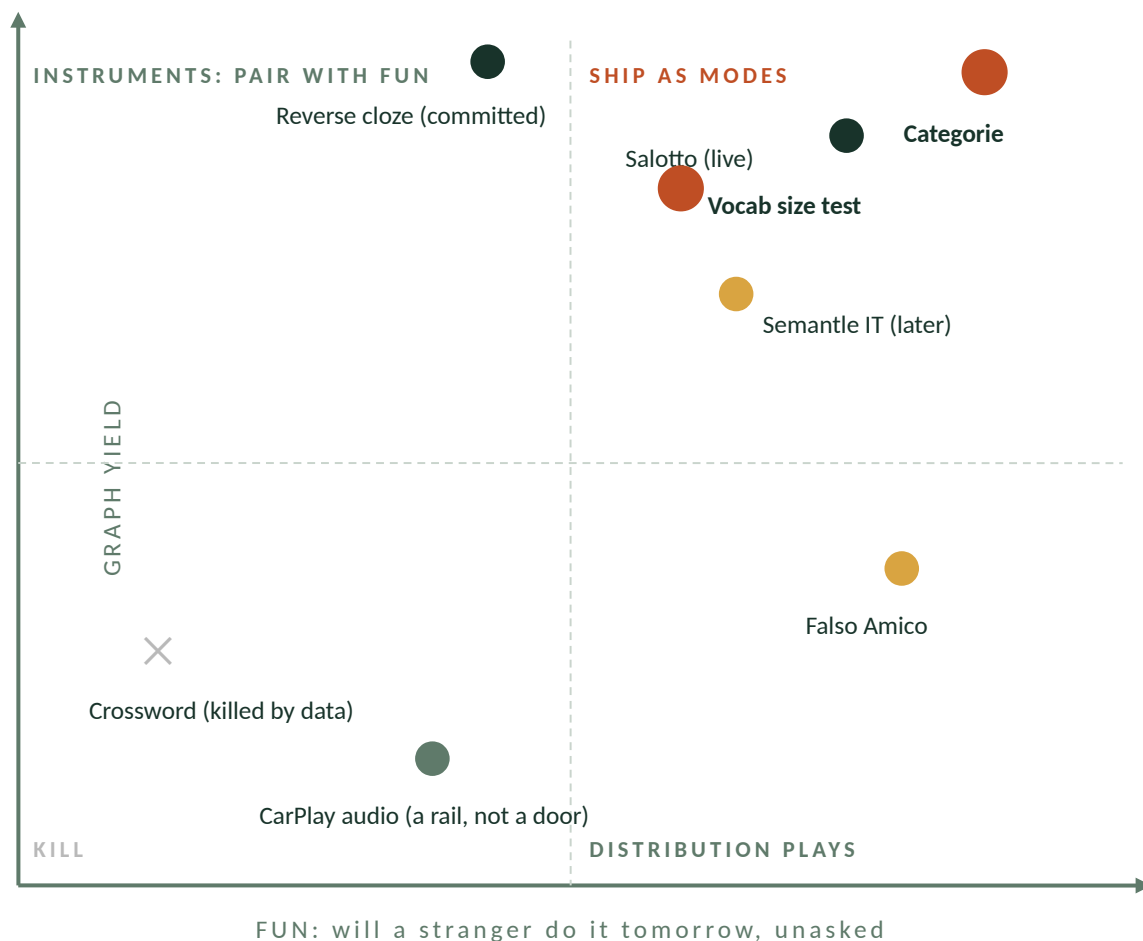
The loop already works when fueled: 30 arrivals from player shares, converting at 47%. The plan fuels it and activates channels one at a time, never two at once, so arrivals stay attributable. Every channel gets a threshold written before it runs.

STEP 0, THIS WEEK per-artifact ref codes on every share link (s=grid, s=test, s=reddit...), so the scoreboard reads arrivals by artifact and by channel.

Channel	When	The move	Threshold, written first
In-product share loops	Always on	Every door ships its artifact: the result grid today, the Categorie grid, the Vocabolario number card next	Share arrivals per week doubles from the ~2-per-week baseline
Reddit	Weeks 1-2	Execute the drafted playbook: r/italianlearning and r/languagelearning posts, modmail drafts already written	50 tagged arrivals; 15% of them produce a word
App Store search	Weeks 2-4	ASO pass on "wordle italiano" and "giochi di parole"; iOS 1.1 is already live	Search impressions to installs, measured in App Store Connect
Vocabolario launch	With its experiment	The score card is the campaign, seeded to Reddit; built for the Stranded A2 who wants proof	25% of finishers share the number; tagged arrivals produce at 20%
WhatsApp group	Demoted to lab	Recruiting ground for Salotto sessions and betas, not a growth channel: 46-to-5 taught that	None; it is no longer asked to grow anything

One channel at a time. Thresholds pre-registered. A channel that misses gets iterated once, then cut, exactly like a door.

Every candidate door, ranked: fun against graph yield



THE MODE-VERSUS-APP RULE

A new door ships as a mode inside Paroletta unless it fails one of three tests: a different cadence (Salotto is scheduled and social), a different identity promise (Filotto is named and accumulating, Paroletta is anonymous), or a different platform capability (background audio needs native). App Store cold start is paid per app; modes share retention and one graph.

TWO HONEST DOTS

CarPlay plots low on yield because passive audio measures exposure, not knowledge: it is a rail that spends the graph, built when the reader is worth carrying. And the crossword sits in the kill quadrant because it already died in production: zero completions, ever. The framework is not post-hoc.

The next doors, as pre-registered experiments

Every launch writes its success metric down before it ships and measures for a fixed window, the discipline the Salotto Marks already proved. Thresholds below are pre-registered proposals: tuned before launch, never after.

Experiment	Success metric, written down first	Share artifact (ref code)	Window
Categorie: 16 words, 4 hidden groups, inside Paroletta	At least 8 words measured per player per day in the mode, and the grid out-shares the word game	Emoji grid of the four solved groups (s=grid)	3 weeks
Vocabolario: the vocab-size test as the front door	40% of starters finish; 25% of finishers share; finishers start with a calibrated graph	The number card: "~3.400 parole · top 8%" (s=test)	2 weeks
Falso Amico: rapid-fire false friends	At least one new producing learner per week arrives from its shares	The gotcha card: "no, morbido non significa morbid" (s=falso)	4 weeks
Filotto reader: the nightly chapter starts sending	60% of sent chapters opened, sustained across two full weeks	The streak card: "il mio filotto: 14 sere" (s=filotto)	2 weeks
Salotto: two community sessions on the record	200 or more distinct lemmas per person per session; scorecard filled in	Session recap: minutes spoken, words used, invite (s=salotto)	1 month

Also in the window: the visible farm (your word count, growing: the retention surface), the opt-in identity seam, and the fixed iOS input flow. CarPlay waits for the reader; Semantle-in-Italian waits for Categorie.

Build the asset now. The graph opens four businesses later.



The reader subscription

Filotto's nightly chapter is a book only you can read, generated inside your measured vocabulary. Consumer subscription, the natural first revenue.



The graph as an API

Tutors, schools, and other apps personalize against a learner's real lexicon. We become the measurement layer for the category that skipped it.



Cohort content

Learners cluster. "A lexicon like mine" powers group readers, club editions, and community material at near-zero marginal cost.



Adjacent languages

Nothing in the machine is Italian. The corpus pipeline, drill, and generation loop are language-agnostic. Italian is the first instance, chosen because we live it.

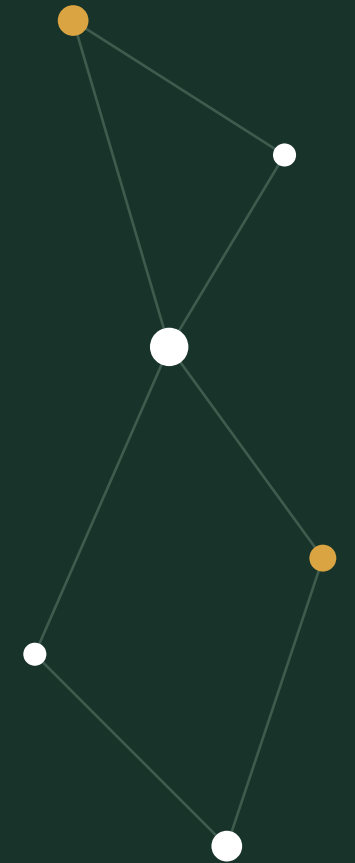
WHAT A RAISE WOULD HAVE TO PROVE the graph growing through every door, week over week; one door acquiring learners at flat or falling cost; one spend surface converting graph into retention. The scoreboard on slide 13 is that proof, or it is not.

Duolingo built the world's best streak. **We are building the learner's lexicon.**

Content is free. **Measurement compounds.** Speed decides who gets there.

The ask depends on who you are. Investor: watch the scoreboard with us for a quarter. Builder: pick a door and own it end to end.
Executive: borrow the method; none of it is specific to language.

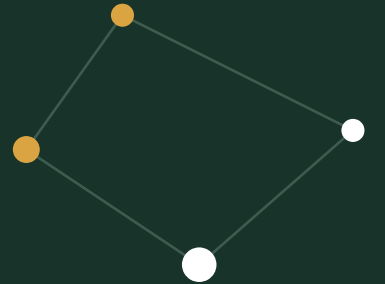
paroletta.com · filottolang.com · salottolang.com



APPENDIX

The operating record

The OKRs as adopted Aug 20, 2026, and the execution plan they kicked off. Kept for posterity: this is what "decide in the morning, ship in the afternoon" looks like on paper.



Aug 20 to Nov 20: three objectives, measured baselines

Cadence: ship day Aug 20 (snapshot, farm, ref codes, iOS build all same-day) • baselines Aug 21 • first grade and recalibration at the Sept 1 pulse • final grade Nov 20 • reviewed daily in the pulse against a live PostHog scoreboard.



O1 • Prove the graph compounds

- True graph size is a chart: snapshot ships day one, baseline next morning
- 50 learners holding 10+ produced words (baseline: 31 producers all-time)
- Weekly produced-word events at 3x the Aug 21 baseline
- The graph gets spent: nightly reader at 60% opens; Salotto sessions at 200+ lemmas/person



O2 • Fix the two funnel leaks

- Touched-to-producing 19% → 40% (post-Aug-21 cohorts)
- Week-two return 4% → 15%
- The visible farm ships day one; 60% of producing learners have seen theirs by Sept 1



O3 • Make distribution compound

- Share arrivals 2 → 10 per week, attributable by artifact (baseline conversion already 47%, twice base rate)
- 50 new producing learners in the window (vs 31 in all history)
- Categorie and Vocabolario launched pre-registered, verdicts written inside their windows
- One external channel passes its threshold and goes always-on

Deliberately absent: revenue (deferred by strategy), DAU (downstream of O2), CarPlay and Semantle and Falso Amico (sequenced behind the verdicts). Grading convention: 0.7 average is a strong quarter.

Five parallel lanes, one scheduler: the builder's hours

A • Paroletta web

Ship day: ref codes, graph snapshot, visible farm, all live same day. Then Categorie (content spike Aug 22, launch week of Aug 31, 3-week window), then Vocabolario (mid-Sept, 2-week window).

B • Paroletta iOS

Input-flow fixes plus the starts-inflation fix, submitted same day. App Store acceptance (est. Aug 25 to 29) is iOS cohort day zero for the O2 conversions.

C • Filotto reader

Send path live this week; first nightly chapters by Aug 31 so the 60%-opens KR has two full weeks of data before mid-Sept. Open-rate instrument decided before first send.

D • Salotto

Session 1 with the community by Aug 28 (invite sent day one), session 2 in September. Transcripts are the richest graph feed: 200 to 400 lemmas per person per session.

E • Channels

One at a time, gated on ref codes: Reddit posts Aug 22 to 25 (verdict Sept 8), ASO weeks 2 to 4, the Vocabolario number card as its own campaign at launch.

CRITICAL PATH App Store review is the only uncontrollable dependency, and it gates only the iOS cohorts. Everything else was designed to be shippable the day it was decided.