

Clinexus — the working papers

Everything behind the estimate, laid out so you can read it at your own pace rather than on a call. We built a model of your product from your Product Overview, your answers of 28 July and your V1.0 designs, and read the three against each other. **Where they disagree, this is where we say so.**

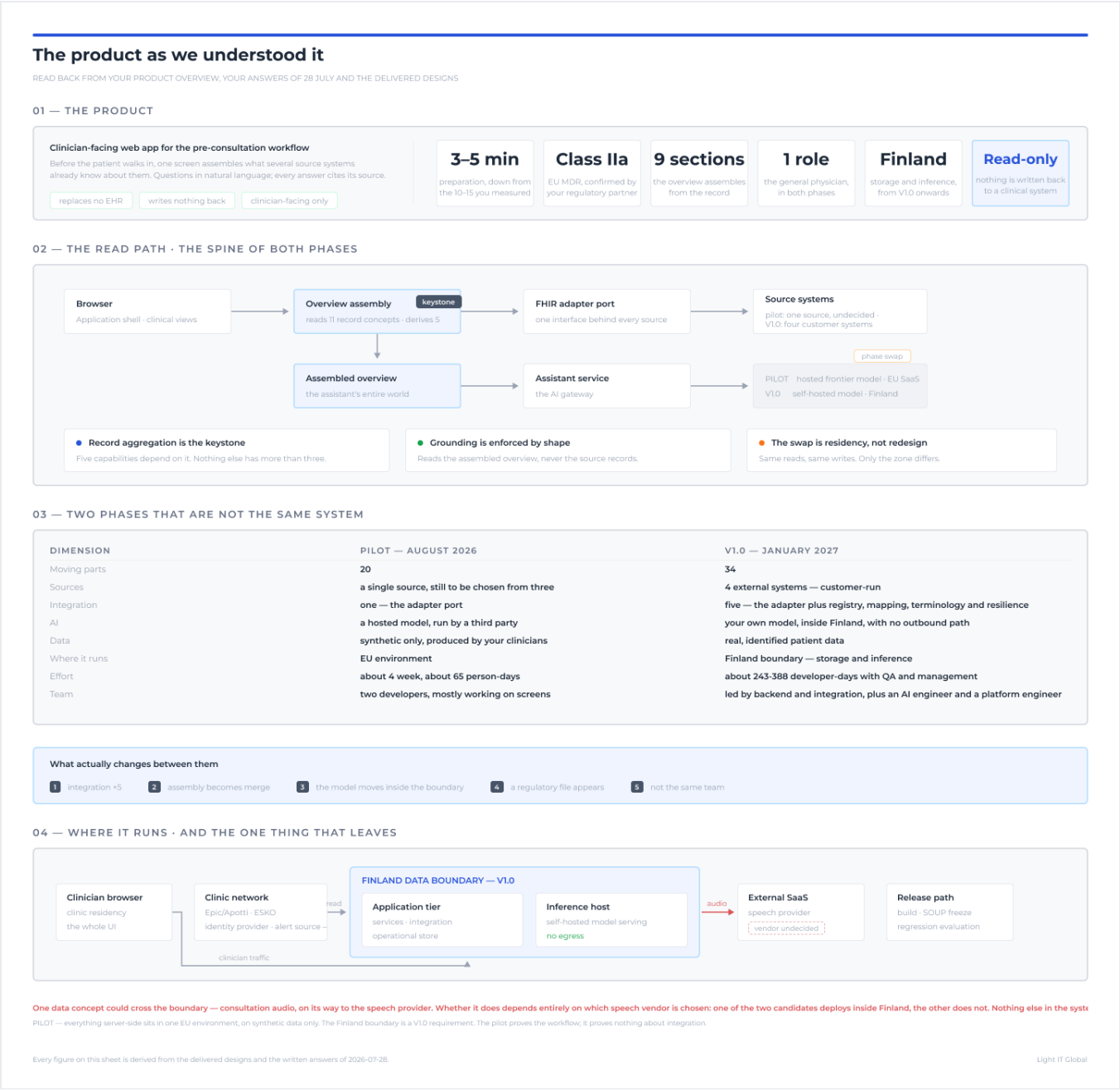
Start here

Two sheets that say what we think we are building. If anything here is wrong, everything downstream is wrong with it — so this is the pair worth arguing with first.

The product as we understood it

Our reading, in one page: what the product is, the path the data takes, the two phases, and where each phase runs.

BOTH PHASES



What actually changes between them

1

 integration ×5

2

 assembly becomes merge

3

 the model moves inside the boundary

4

 a regulatory file appears

5

 not the same team

Clinician browser

clinic residency
the whole UI

Clinic network

Epic/Apotti · ESKO
identity provider · alert source

FINLAND DATA BOUNDARY — V1.0

Application tier

services · integration
operational store

read

Inference host

self-hosted model serving
no egress

External SaaS

speech provider
vendor undecided

audio

Release path

build · SOUP freeze
regression evaluation

One data concept could cross the boundary — consultation audio, on its way to the speech provider. Whether it does depends entirely on which speech vendor is chosen: one of the two candidates deploys inside Finland, the other does not. Nothing else in the system.

PILOT — everything server-side sits in one EU environment, on synthetic data only. The Finland boundary is a V1.0 requirement. The pilot proves the workflow; it proves nothing about integration.

Every figure on this sheet is derived from the delivered designs and the written answers of 2026-07-28.

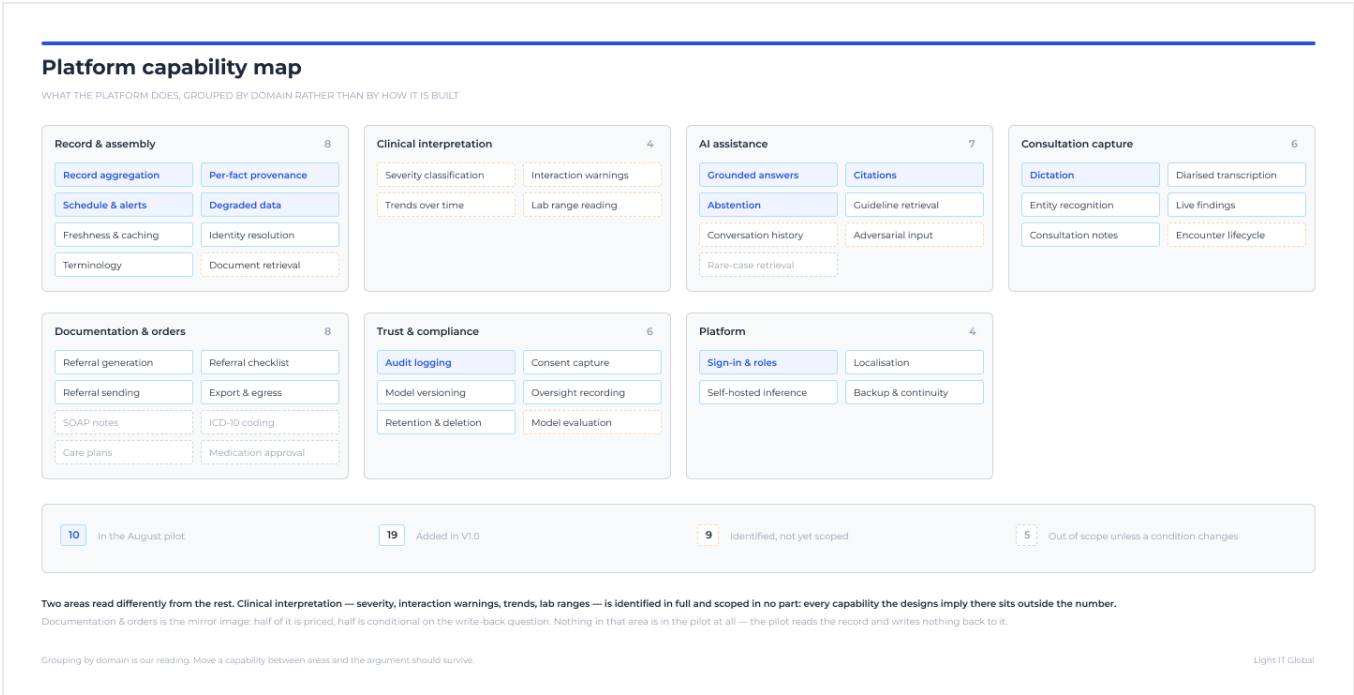
Light IT Global

Read back from your Product Overview, your 28 July answers and the V1.0 designs.

Platform capability map

What the platform does, grouped by domain rather than by how it is built — shaded by phase and by whether it is inside the estimate.

BOTH PHASES



Grouping by domain is our reading. Move a capability between areas and the argument should survive.

The pilot

The August pilot is the part we can start on now. Nothing in it waits on access to a system somebody else controls, which is what makes the date reachable.

Architecture & deployment — pilot and V1.0 side by side

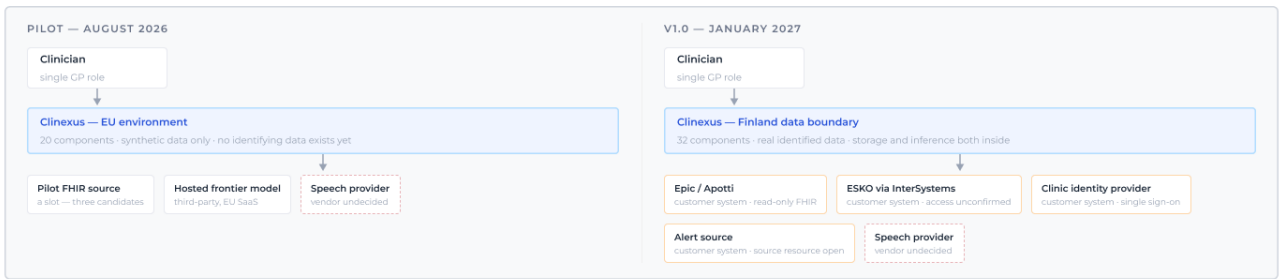
Three zoom levels, both phases next to each other, plus the container reference. It sits in this section because the pilot is the left-hand column — but read the whole sheet: the point of drawing the phases together is the difference between them.

BOTH PHASES

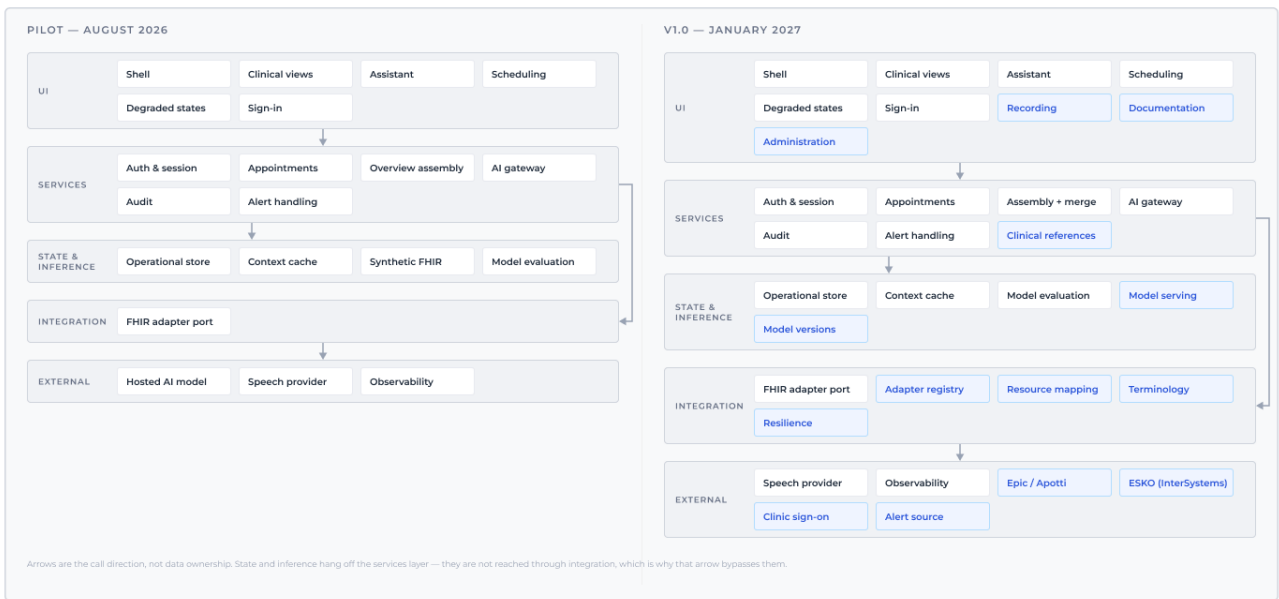
Architecture & deployment — pilot against V1.0

THREE ZOOM LEVELS, THE SAME TWO PHASES SIDE BY SIDE - MOST OF THE PILOT SURVIVES INTO V1.0 - 2026-07-30

A — L1 · CONTEXT · WHAT THE SYSTEM TOUCHES



B — L2 · CONTAINERS · WHAT IT IS MADE OF



B2 — WHAT EACH CONTAINER IS, AND WHOSE IT IS

THE INTERFACE — and what each part of it draws

Capabilities are realised in the services behind these, so the interface is described by what it renders instead — regions, screens, and the screens the client has not drawn.

Web app shell both

Routing, the navigation rail, the patient banner and the tab shell. Five regions covering all 46 screens between them — only the rail is on every one — change one and everything retests.

Assistant UI both

Chat, pre-screening, starter prompts and the citations under each answer. Fourteen regions across 42 screens, plus a clinical reference browser that does not exist yet.

Recording UI V1.0

Transport controls, the live transcript and the highlights drawn over it while the consultation runs. Eight regions across nineteen screens.

Degraded and error UI both

Missing, outdated and conflicting values labelled in place. Neither state exists in the delivered designs — both are ours to design as well as build.

Clinical data views both

The clinical surface the doctor reads: overview, history, prescriptions, rare cases. Nineteen regions across 44 screens, plus two we must add — a document viewer and a note editor.

Documentation UI V1.0

SOAP editor, care plan and referral composer with its send gate. Ten regions in only two screens — the densest part of the product. Consent capture and an export view are still missing.

Scheduling UI both

The day's appointment list and the visit detail behind each row. Six regions, two screens.

Auth UI both

Sign-in and single sign-on. Both screens are ours: the delivered designs include no sign-in and no session-expiry state.

Administration UI V1.0

Five operator surfaces that appear in no design at all: audit log viewer, section-to-source configuration, model versions, retention, clinic users.

WHAT WE BUILD BEHIND IT

Services, state, integration and inference. This is where the capabilities live, and where almost all of the V1.0 growth happens.

Overview assembly service both

Builds the one-screen overview: reads eleven record concepts and derives five of its own. In the pilot it assembles from one source; in V1.0 it also merges across several with precedence rules. Carries record aggregation, identity resolution, trends, lab ranges and severity.

FHIR adapter port both

The one internal interface behind every source; twelve data concepts pass through it. It is what makes connecting several systems bounded work rather than open-ended.

Session and context cache both

Holds the assembled overview between requests so the record is not re-read on every question. Where the freshness policy lives.

Appointments service both

Reads the day list and the alert tiles from the source. Computes neither.

Alert and flag handling both

Ingests alerts from the source and routes them to the surfaces that display them.

Adapter registry V1.0

Which adapter serves which source, once there is more than one.

Terminology service V1.0

Code systems and units normalised so that two sources can be read as one.

Model serving layer V1.0

Self-hosted inference inside Finland. Reads and writes exactly what the hosted pilot model does — only the place changes.

Regression evaluation suite both

The standing test set that stands in for the anomaly list no model vendor publishes.

WHAT WE DEPEND ON — the parts that are not our work

Four are the client's own systems, two are products we buy, one is an empty slot. Every one of them can fail, change or refuse access without us touching a line of code.

Epic / Apotti customer system V1.0

The customer EHR, read-only over FHIR. Six clinical concepts come from here.

Clinic identity provider customer system V1.0

The customer identity provider. We map claims to the physician role and keep no directory of our own.

Hosted frontier model bought pilot

Pilot only. Replaced by self-hosted serving in V1.0 — same reads and writes, different place.

Speech provider undecided both

A slot, not a product. Where it runs follows the vendor choice — and so does whether any data leaves Finland at all. It is the only component that decides that.

Assistant service both

The AI gateway. Builds each prompt from the assembled overview — never from the source records — and enforces grounding, citation and abstention. Ten capabilities, more than any other container.

Operational store both

PostgreSQL. Everything the product genuinely owns — sessions, audit events, conversations, generation records, oversight decisions. Nineteen concepts.

Auth and session service both

Sign-in, session lifetime and enforcement of the single physician role.

Audit service both

Records authentication, patient access, generated overviews and queries. An emitter in the pilot; an emitter with a read path in V1.0.

Clinical references service V1.0

Approved external clinical material, kept visibly separate from patient data.

Resource mapping V1.0

Which FHIR resource fills which overview section, and which source wins when two disagree.

Resilience layer V1.0

Retry, circuit-break and degrade rather than fail when a source is slow or absent.

Model version records V1.0

An immutable record of which model and prompt produced each output. Technical-file evidence, not a configuration table.

Pilot FHIR source pilot

The pilot reads from one source and it has not been chosen. Our own synthetic server is priced; the public Epic sandbox and the LAMK profiles are the other two candidates named in the answers.

ESKO via InterSystems customer system V1.0

The Oulu-region EHR. Access is unconfirmed and the interface surface is unknown to us — the widest unknown in the programme.

Alert source customer system V1.0

Whatever system carries an alert on the EHR side. Which resource that is remains an open question.

Observability stack bought both

Logs, metrics and traces.

C — L3 · DEPLOYMENT · WHERE IT RUNS

PILOT — AUGUST 2026

Clinician browser clinic · 6 components

shell · clinical views · assistant · scheduling · degraded states · sign-in

PILOT ENVIRONMENT EU · 11 components · synthetic data only

Everything server-side, in one box

auth & session · appointments · overview assembly · AI gateway · audit · alert handling
adapter port · operational store · context cache · synthetic FHIR server · observability

no clinic network nothing to integrate with yet

the one source is ours, synthetic, and always answers

External SaaS external · 2

hosted AI model · speech provider

Release path external · 1

regression evaluation

The pilot's entire server side is one box in the EU. V1.0 splits it into a Finland boundary with an isolated inference host — and the clinic network, absent from the pilot, appears with four systems we do not control.

V1.0 — JANUARY 2027

Clinician browser clinic · 9 components

+ recording · documentation · administration

FINLAND DATA BOUNDARY storage and inference both inside — the residency requirement made physical

Application tier finland · 15

services · integration · store · cache · observability

Inference host finland · 2

model serving · version records · no egress

Clinic network clinic · 4 systems

Epic/Apotti · ESKO · identity provider · alert source — all customer-run

External SaaS external · 1

speech provider only

Release path external · 1

regression evaluation

Both phases are drawn from the same model, so a part shown in one column and absent from the other is genuinely absent.

Light IT Global

The clinical seam — assistant plus model — sits behind one interface from day one.

Pilot navigation

How a doctor moves through the pilot, and the frame that persists around every screen. Solid arrows are drawn in your designs; dashed amber ones are our reading.

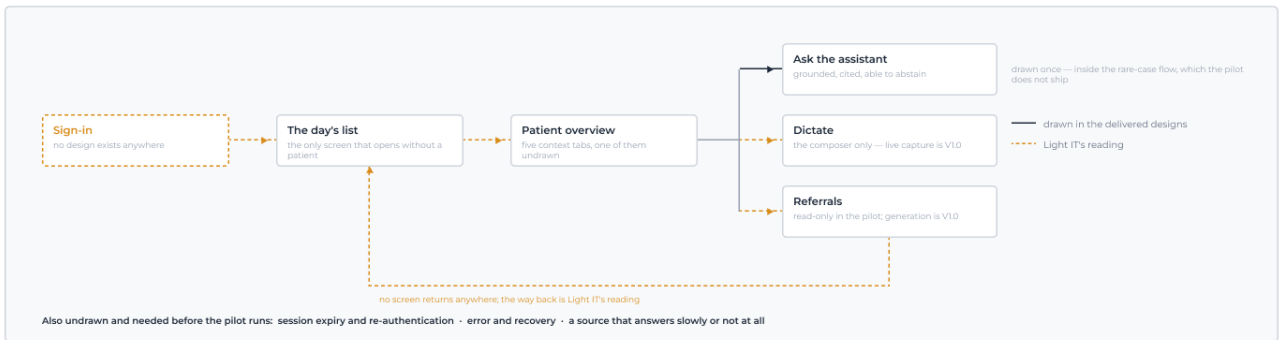
PILOT

How a doctor moves through the pilot

THE PATH AS LIGHT IT READS IT - DRAWN TRANSITIONS AGAINST INFERRED ONES

A — THE PATH

The delivered designs contain six families of screen and two transitions between families in the whole file. Every hop below except one is Light IT's reading of how they join up.



B — THE FRAME AROUND EVERY SCREEN

These parts persist across the screens. Three of them have no written requirement behind them at all, and two have no design for the state they have to reach, so the last column is where Light IT has committed to a reading.

PART OF THE FRAME	WHERE IT APPEARS	WHAT THE DESIGN LEAVES OPEN	WHAT LIGHT IT READS IT AS
App rail	every screen	nine destinations, only patients is ever active	eight are V1.0 promises — the pilot renders one and does not show the rest as dead controls
Breadcrumb	every screen once a patient is open	only the third crumb ever changes	it names the patient and the open tab, and nothing else
Patient banner	every screen once a patient is open	six vitals with no recency, no source and no alarm styling	every value carries when it was taken and where it came from, and out of range is styled as out of range
Context tabs	patient overview only	five tabs, one of which has no screen on either side	four tabs ship; the fifth stays out until a design exists
Co-pilot drawer	patient overview and pre-screening	no written requirement anywhere	not a second assistant — the same one, docked
Ask AI launcher	patient overview and live capture	a floating control with no destination drawn	opens that same assistant in that same session, so a question asked in one place is visible in the other
Note dock	every screen once a patient is open	collapsed in every frame; the open state was never designed	one dock holding three kinds of note, and its open state is a pilot screen Light IT has to design

The frame is the product's only continuity, and it is the part the designs specify least.

A doctor never sees a screen — they see this frame with something inside it. Three of the seven parts rest on no requirement at all, which is why each one above carries a stated reading rather than a question.

Screen families, transitions and frame parts are read from the delivered designs; every inferred hop is recorded in the model and can be corrected against a single source.

Light IT Global

On the pilot path exactly one arrow is drawn in the delivered file. The rest is our reconstruction — said on the sheet rather than presented as fact.

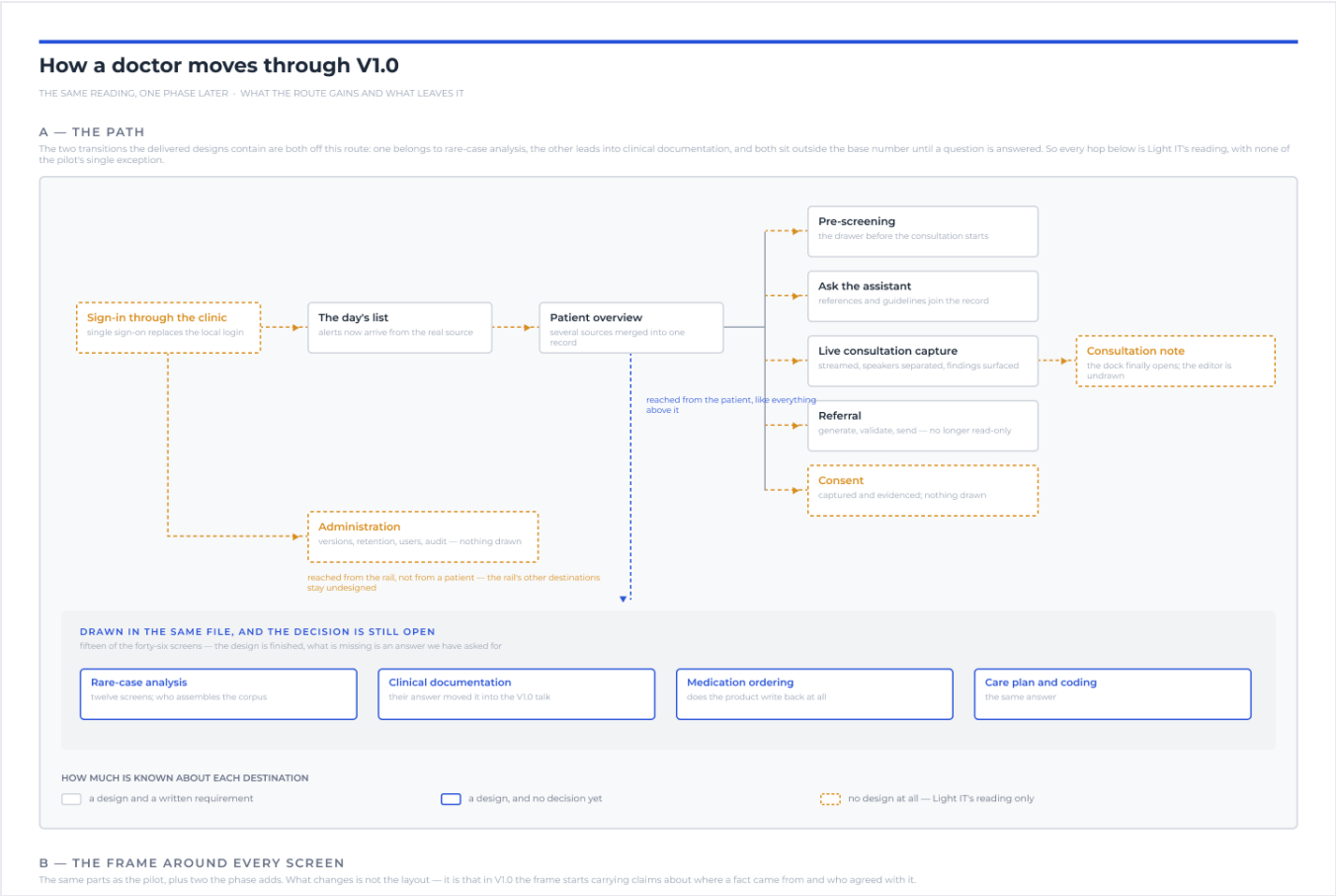
Clinexus V1.0

V1.0 is the same screens over materially more depth. Most of what gets added renders nothing — which is the honest answer to **why V1.0 costs several times the pilot while looking much the same**.

V1.0 navigation

The same route one phase later, drawn to the same shape so the two can be read down the page. Its finding is sharper than the pilot's: on the V1.0 route, none of the hops is drawn in the delivered file.

V1.0



PART OF THE FRAME	WHAT V1.0 CHANGES	WHAT THE DESIGN LEAVES OPEN	WHAT LIGHT IT READS IT AS
App rail	administration makes settings a live destination, so the rail stops being a list with one working entry	six of the nine still have no screen behind them	patients and settings ship; medications follows the band below; the other six stay out and are not rendered as dead controls
Breadcrumb	new third crumbs — live capture, documentation, referral letter	It never names which system a fact came from	It names the patient and the open context and nothing else; source belongs in the banner, not here
Patient banner	the patient stops being one record — several systems are merged behind it	no timestamp, no source, no alarm styling, and nothing at all for a disagreement between systems	this is the frame's real V1.0 change: each value carries when and from where, and where two systems disagree the banner shows both rather than picking. A wrong merge is the worst failure this product has
Context tabs	two of the five tabs wait on decisions, one has no design anywhere	In Depth Study has no artboard on either side	Overview and Patient History ship; Rare Case and Diagnosis follow the band below; In Depth Study stays out until a screen exists
Co-pilot drawer	pre-screening is in scope, so the drawer becomes real work rather than an undocumented container	Its Reports card links out to documents nobody has designed a viewer for	the same assistant, docked — and the Reports card is blocked on the document viewer in the list below
Ask AI launcher	the assistant now also reaches clinical references and guidelines	still no destination drawn for the floating control	one conversation and one session, whichever of the three entry points opened it
Note dock	it finally opens — consultation notes ship in V1.0	the open state was never designed, and it is labelled Note on some screens and Notes on others	one dock, three kinds of note, and the editor is the first of the asks below
Language switch	new — localisation is in V1.0	no language is ever named anywhere, and no control exists	it belongs in the frame beside the user avatar, not inside a settings screen a clinician has to leave the patient to reach
Accept or override	new — every AI-produced element needs a recorded human decision	nothing in any screen shows who agreed with an AI output	this lives in the frame rather than on one screen: wherever the product generates, the same affordance records who accepted or overrode it, because that record is what the device file rests on
Two of the nine are new, and both are claims rather than controls. A language switch and an accept-or-override affordance are not screens anyone asked for — they are what a device has to show on every screen once it generates text and once it merges records from more than one system. Putting them in the frame is cheaper than discovering later that each screen needed its own.			
C — THE DESIGNS V1.0 NEEDS AND DOES NOT HAVE These are the destinations V1.0 has to ship with nothing drawn behind them. Each is written up from a requirement instead of a screen, which is where the width of the estimate comes from — and each is a design request rather than a question.			
The consultation note editor the dock also collapsed on five screen families and was never opened Consent capture and its record required before anything is captured at all Retention and deletion the erase-this-patient screen Model and prompt version administration which version answered, and the ability to pin it Audit log viewer written in the pilot and never read Section-to-source configuration which record section reads which resource Clinic user administration Clinical reference browser Document viewer Export and print view Language selection A way to reach a patient who is not on today's list the product opens on the day's list and no screen anywhere goes further Imaging summaries named in the record and given no surface — show them, link out, or say they are out These are asks, not risks. Everything on the diagram above in amber is here, plus two destinations the route never reaches at all. Getting a screen for each is the single cheapest way to narrow the V1.0 range, because it replaces a written assumption with something that can be looked at and disagreed with.			
Screen families, transitions and scope states are read from one model; the fifteen-screen count is a traversal, not an eyeball.			
Light IT Global			

Panel C lists the destinations V1.0 must ship with nothing drawn — a screen for each is the cheapest way to narrow the range.

The AI — questions, model classes, providers

What actually goes to a model, the classes of question it has to serve, the licence filter that removes most candidates before any evaluation, and the named ones that survive it.

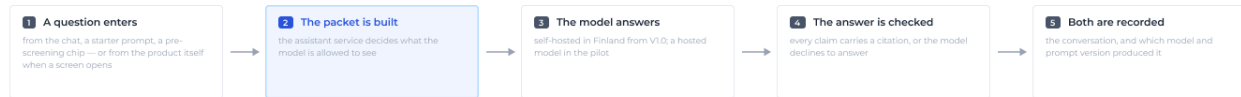
V1.0

The AI — what it answers, what runs it, who supplies it

Self-hosting in Finland was decided in one sentence. This is what follows from it.

01 — HOW A QUESTION REACHES A MODEL, AND WHAT GOES WITH IT

The same path serves every question the product asks. The only thing that changes is how much of the record travels with it — and that is not decided.



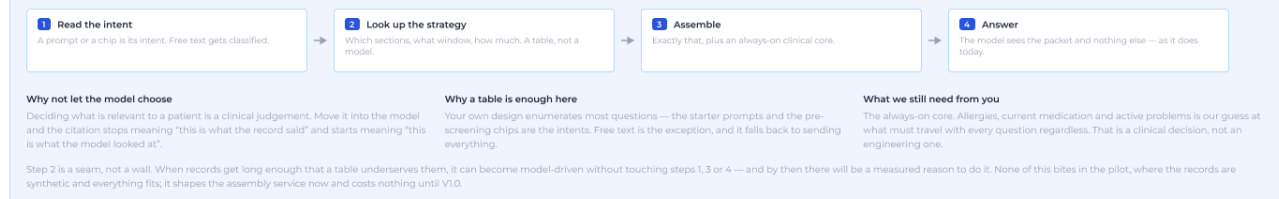
The rules answer only from what follows; cite every claim; decline rather than guess	Part of the record the assembled overview — and how much of it is the open question below	The question typed, or the text behind a prompt or a chip	Retrieved material a guideline or a similar case, when the question calls for one
--	---	---	---

Who decides what the model may see — the choice that matters most, and the one not yet made

A record grows; a context window does not. Whatever is left out, the model cannot know — and it answers anyway. The real axis is not how much travels, but who picks it.

Send everything, every time Deterministic and trivially evidenced — the answer saw the whole record. Fails on record length, not on logic.	A strategy per question type Each kind of question maps to the sections it needs. Testable, versioned, and a clinician can read it. Free text is the exception it has to fall back from.	Search the record per question Scales to any record. A miss is silent, and a missed section produces a confident answer on an incomplete one.	Let the model fetch what it wants Most capable on a long record and a complex question. But deciding what is relevant to a patient is itself a clinical judgement, and a multi-step fetch loop is hard to reproduce.
--	--	---	--

Our position — decide in stages, so that the model never decides



02 — That is four classes of model, not one

It reads as one decision about one model. It is four, and only two concern a language model.

Speech recognition Deciding what is relevant to a patient is a clinical judgement. Move it into the model and the citation stops meaning "this is what the record said" and starts meaning "this is what the model looked at".	Generative language model Answers, summaries, notes, letters. What the Finland decision is really about.	Retrieval and embedding Cases and guidelines. Cheaper to run, far easier to evidence — it returns documents, not sentences.	Extraction Entities in running text. Constrained prompt or a small dedicated model — a real cost difference, still open.
--	--	---	--

03 — What each model is allowed to see

The safety property is not what the model is — it is what the model is given.

The language model never sees the source records It reads the overview the platform already built and can cite. A prompt can be argued with; a topology cannot.	Consultation audio is the only thing that could leave Finland Everything else stays inside under every option. Audio is the single conditional: one speech vendor deploys on your own infrastructure, the other processes in the EU. Choose the first and nothing leaves at all.
Into the model The overview, the question, and a retrieved guideline or case. Nothing else.	Out of the model The answer, its citations, and which model and prompt produced it.
Kept Conversations, generated artefacts, generation records. Inside Finland, ours.	Never kept Consultation audio. It passes through and lands nowhere.

03b — What the record can smuggle into a prompt

A prompt is assembled out of the record. A record is text other people wrote, and it can contain a sentence shaped like an instruction.

Record content is untrusted input, always A referral letter, a free-text note, a comment on a medication line — any of it can carry wording that reads as a command. It arrives inside the same prompt as ours, and nothing in the text marks which is which.	What it could do Change the answer, make the assistant drop its citation rule, or make it speak about a patient nobody asked about	What defends it Keeping record content separate from instructions, refusing to act on anything that arrives as data, and testing that the refusal holds under pressure	Where it stands today Absorbed. Light IT would build it and it appears in no line of the estimate — a choice, written here so that it is a visible one
---	--	--	--

04 — RARE CASES, AND HOW THEY WOULD ACTUALLY WORK

The most consequential feature in the file and the least described. It is also the only one that argues with itself: the rarer a case is, the more useful it is to a clinician — and the easier it is to identify the patient in it.



Structured match, not embedding search

why

A shared diagnosis code is a reason you can put on the screen. A similarity score is not. For a device that has to explain why it showed you something, the explainable method wins even where it retrieves less.

Suppress cases with no neighbours

the safeguard

Rarity is the identifier. Removing the name from a case that is one of a kind removes nothing. So: show a case only if enough similar ones exist — which means the feature deliberately withholds its most striking material.

Retrieval, not a verdict

the boundary

The product returns cases and says why each matched. The clinician compares them. The recommendations and prognosis your screen shows are the part we would not build — that is a conclusion, and it belongs to the doctor.

Three things this design cannot settle, and none of them is technical

Where the corpus comes from

Your own clinic records, a licensed case collection, or a curated set. Each answer has a different legal route and a different owner, and none of them exists today.

In the pilot this is demonstrable and safe: the corpus is the synthetic patients your clinicians produce, so the mechanism can be shown working without any of the above being settled. That is the sequence we would recommend — prove the retrieval on synthetic cases in August, and settle the corpus and its legal basis before it touches a real one.

What makes a case rare

A diagnosis on a rare-disease list is a lookup. An unusual combination of findings is a statistic computed over the corpus. Different work, different cost, and your design does not say which.

Whether it is lawful at all

Comparing one patient against others is a secondary use of health data. Under the Finnish rules that needs a basis, and your regulatory partner should establish it before any of this is built.

05 — WHAT RUNS THE AI, AND WHAT THE FINLAND DECISION LEAVES

The two phases are answered differently, and the pilot does not need a model chosen at all.

Pilot — a hosted model

August 2026

The data is synthetic, so residency does not yet apply. Best available capability, no infrastructure, nothing to select. The pilot needs no model decision.

V1.0 — your own model, in Finland

January 2027

One sentence in your answers removes every model that exists only as a hosted service — the whole frontier tier. What is left is open weights, and that is where the real filtering starts.

Then the licence removes most of what is left

“...the unauthorized or unlicensed practice of any profession including medical or health” — Llama. “...automated decisions in domains that affect well-being, including healthcare” — Gemma 3 and earlier, which also reserves the right to switch you off remotely. The medical-specialist model, MedGemma, needs device approval before you may use it at all.

What survives is small and clean

EuroLLM · Qwen 3.x · Poro-34B · Gemma 4 since April 2026 — all Apache 2.0, no medical clause, nothing to argue in the technical file. Add the hardware ceiling of one clinic on one card and the field is four names before anything is downloaded.

And then it is measured, not read

Faithfulness to the record, citations that hold, and Finnish — tested on Finnish clinical scenarios built from your own synthetic patients. No public benchmark measures any of those three.

Worth knowing regardless of which name wins: open weights are not an open licence, and nobody on this deal had checked. It removes more candidates than the hardware does.

06 — The language models that survive, and where they come from

Nothing is chosen. These are what reach the harness.

Poro 2 Instruct — 8B

Finland · LumiOpen

AMD Silo AI, TurkuNLP, trained on LUMI>About +10 points on Finnish for -1 on English. Carries the Llama licence flag.

Poro-34B

Finland · LumiOpen

The only candidate both Finnish-native and licence-clean, 2024 generation — a reference point, not a favourite.

EuroLLM-Instruct — 9B and 22B

EU consortium

All 24 EU languages, Finnish and Swedish named. Dec 2025, Apache 2.0. The only candidate with no asterisk — and a consortium, not a vendor.

Gemma 4

Google · open weights

The capability ceiling that still fits one card. Finnish is covered, not targeted — exactly what the harness is for.

Qwen 3.5 / 3.6

Alibaba · open weights

Longest context of the group. Chinese origin in a Finnish public-health supply chain is a procurement conversation — better raised than discovered.

OpenEuroLLM 8B

EU · expected

Same consortium, expected mid-2026. A watch item.

07 — Speech is a separate market, and it is where the Finland question was actually decided

Medical Finnish is the hard part. General models are 10-12 % wrong on ordinary Finnish, worse on clinical words.

Speechmatics

UK · deployable on premise

Has a Finnish medical model, and ships as containers you run yourself — so medical accuracy and strict residency at once. No Finnish error figure is published; the Nordic ones that are run 3.9 to 8.0.

Inscripta Oy

Finland · healthcare specialist

Finnish healthcare specialist, named county customers — the better market story. EU processing; on-premise not advertised.

A general model we host ourselves

any origin

Passes residency, fails accuracy. Closing the clinical gap is a research programme, not an integration.

08 — What we would actually run, and what it costs to find out

Four on the harness. The free filters cut the field first — that is what keeps it to days.

The shortlist

two to four days

EuroLLM 22B baseline · Gemma 4 ceiling · Poro 2 8B Finnish reference · Qwen for context. Whether a general 22B beats a Finnish-tuned 8B is the most useful number it can produce.

One hosted frontier model, on synthetic data only

not a candidate

A ceiling during the pilot, on synthetic data where residency does not yet apply. Turns “self-hosting costs you capability” into a number.

09 — How anyone would know the model was wrong

Every other third-party component arrives with a published list of its known faults. A language model does not, and will not.

There is no anomaly list to read, so one has to be built

IEC 62304 expects the supplier’s known-anomaly list for any component we did not write. For a language model no such list exists, so the safety evidence becomes a test set we author and the customer signs.

What the set holds

A fixed body of questions with expected behaviour: does the answer stay inside the record, does it cite, does it decline when the record cannot support an answer

Who signs it

What the set tests for is a clinical judgement. Whoever approves it is defining what correct means for this product

When it runs

On every model version, every prompt change and every provider switch. Each of those is a change, and the set is what says it is still the same product

What this sheet is not

Not a recommendation. Nothing is selected until it is measured on Finnish clinical scenarios — built from the synthetic patients your clinicians produce. No public benchmark tests faithfulness to a supplied record in Finnish.

Perishable: two of these did not exist a year ago. The filters are durable, the names are not.

Licence terms from each family’s published policy. Where a vendor publishes no figure for Finnish, none is quoted here — and we have measured none of them ourselves.

Light IT Global

Your answer 15 — self-hosted in Finland — is what rules out every hosted-only model.

Data — entities, sources and catalogues

Where every fact on the screen comes from, what hangs off what, the six catalogues behind the pickers, and the three ways a clinician can be authorised against your record systems.

V1.0

Data — what the screens stand on

WHERE EACH FACT COMES FROM · WHAT HANGS OFF WHAT · THE LISTS BEHIND THE PICKERS

A — WHERE EVERY FACT ON THE SCREEN COMES FROM

The chain from a standard interface to a pixel, for each thing the overview shows. The structure is the same everywhere — the code systems are not, and that is where the integration cost sits.

WHAT THE DOCTOR SEES	THE THING BEHIND IT	FHIR R4 RESOURCE	CODES — FINLAND	CODES — US CORE
Diagnoses	Diagnosis	Condition	SNOMED CT · Finnish reference sets	ICD-10-CM · SNOMED US
Current medication	Medication statement	MedicationStatement	ATC · VNR nordic article no.	RxNorm
Pending medication	not modelled yet	MedicationRequest	same as above	RxNorm
Allergies	Allergy	AllergyIntolerance	SNOMED CT	SNOMED US · RxNorm
Labs	Lab result	Observation · DiagnosticReport	FinLOINC · Kuntaliitto test codes	LOINC
Vitals banner	Vital sign	Observation	FinLOINC	LOINC
Reports and history	Clinical note	DocumentReference	national document types	US Core note types
Imaging summaries	Imaging summary	ImagingStudy	DICOM · SNOMED CT	DICOM · LOINC
Referrals	Referral	ServiceRequest	SNOMED CT · national	SNOMED US
Encounters	Encounter (source)	Encounter	national encounter classes	US Core class
Appointment list	Appointment	Appointment	national	US Core
Clinical alerts	Alert	undecided — Flag or DetectedIssue	unknown until the resource is named	same
Specialist picker	Specialist directory — source open	Practitioner · PractitionerRole	national registry identifiers	NPI

The structure carries over. The codes do not.

Reading a resource and putting it on a screen is the same work under any profile — FHIR R4 is FHIR R4. Everything that needs the code rather than the label is different: interactions, reference ranges, and deciding two sources mean the same thing.

Which is why the pilot is safe and V1.0 is not.

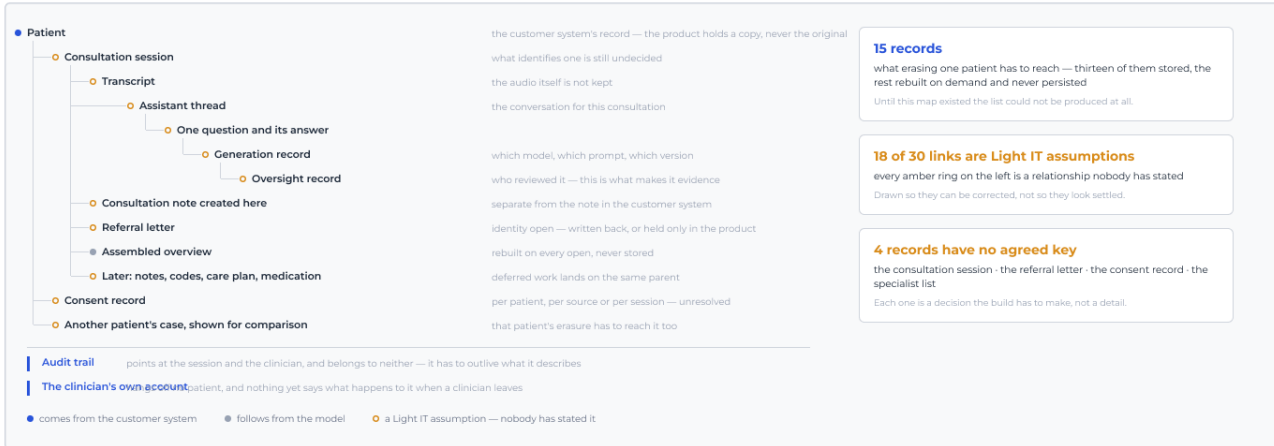
The pilot displays, so it needs the label. V1.0 interprets and merges, so it needs the code. The profile family is nearly free to get wrong in August and expensive to get wrong in January.

Three rows have no source at all.

Pending medication, clinical alerts and the specialist picker. Each is drawn on a screen and each is an open question rather than a mapping task.

B — WHAT HANGS OFF WHAT

Erase a patient and everything on this tree goes with them. The one record that must not be drawn outside it.



C — THE LISTS BEHIND THE PICKERS

Six catalogues the screens need. None is patient data, none comes for free, and a stale one is a clinical problem rather than a cosmetic one.

CATALOGUE	WHERE IT APPEARS	CAN THE CUSTOMER SYSTEM SUPPLY IT	WHO KEEPS IT CURRENT
Care organisations	referral — the institute step	yes, if it exposes Organization	open
Clinical specialties	referral — the specialty step	yes, through PractitionerRole	open
Specialists	referral — who actually receives it	yes, through Practitioner	open — and a specialist who has left is a referral that goes nowhere
Urgency levels	referral — Routine, Urgent, Emergent	partly — the standard carries four codes and none of them is Emergent	Light IT, with a stated mapping
Referral completeness rules	the eleven-point check before sending	no — nothing in the standard says this	Light IT, on customer clinical sign-off
Vital signs shown, and what reads as abnormal	the patient banner	no — the standard names the measurements, not the choice or the thresholds	Light IT, on customer clinical sign-off
Not one of the six has an owner in the plan today — no effort is priced anywhere for keeping them current. Three of them are clinical content Light IT would be authoring. In a Class IIa device, changing what reads as abnormal is a release event, not a setting.			

One model behind every panel here — a record on the tree but missing from the catalogue table is a contradiction, not an omission.

Light IT Global

E — WHAT MAY CROSS THE BOUNDARY, AND WHO SAYS SO

Panel D asked who lets a clinician read the record. This asks the same about data leaving. Three crossings look like one feature — getting a document out — and each answers to a different authority.

THE CROSSING	WHAT ACTUALLY MOVES	WHOSE PERMISSION IT NEEDS
Reading a document from the customer system	the report and letter files the record links out to, fetched on demand	the customer system's, per clinician — the same question Panel D asks about the record itself, and it inherits whatever answer that gets
A clinician exporting from the product	whatever is on screen, into a file that leaves the product and stops being governed by it	nobody has said whether this is allowed at all. It is drawn as a control and has no rule behind it
The product's own backups	everything stored, on a schedule, to wherever the backups live	a residency question rather than a feature one — and the crossing nobody counts as egress until an auditor does
Three surfaces, three permissions, three lines — not one. Priced together they would read as a single integration and be cut as one. Two of the three have no rule behind them today, and the one that does inherits its answer from the sign-in question rather than having its own.		

D — WHO THE DOCTOR IS, TWICE

The customer identity provider says who the person is. The customer record system says which practitioner. The product holds an account and a session. Nothing in the plan maps the three, and which of them carries the link is decided by how a clinician signs in to the record system at all.

HOW A CLINICIAN SIGNS IN TO THE RECORD SYSTEM	WHO DECIDES WHICH PATIENTS THEY MAY SEE	WHERE THE LINK TO THE PRODUCT ACCOUNT LIVES	WHAT THE CUSTOMER AUDIT SHOWS
The clinician authorises the product directly the Light IT assumption, and what is priced	the customer system it keeps the judgement it already makes	inside every token — never stored nothing to keep current	the doctor who looked
One account for the whole product a condition, not a discount	the product the product becomes the access-control authority for patient data	a table somebody maintains a seventh catalogue, and a stale row blames the wrong doctor	the product, never the doctor
Network access and a service user realistic for the older systems reached through a middleware layer	the product same as above, and no protocol-level clinician identity at all	a table somebody maintains plus bespoke integration work	nothing standard
None of this appears in the pilot — the sign-in is local and the record source is synthetic, so there is no second identity to reconcile. That is why the question has to be asked in August rather than discovered in January: the two answers are not a cheap and an expensive version of the same build, they differ in who carries the duty of deciding who may see a patient.			

Erasing one patient has to reach fifteen records. The audit trail sits outside the tree on purpose — it has to outlive both the session and the clinician.

Effort and pricing

Both phases in one place: who does the work, how long it takes, and what it costs. **The pilot is a firm figure; V1.0 is an orientation** — and the difference between those two words is the whole argument of the section after this one.

What the pilot takes — the team, week by week

The staffing behind the August window: hours per role per week across the five weeks, rather than one number. The shape is the part you can check.

PILOT

WHO IS ON IT	HOURS	FTE	W1	W2	W3	W4	W5	WHAT THE LINE IS
Backend and integration	180	0.90	40	40	40	40	20	Carries the analyst by decision — the working sessions with you and the section-to-source mapping ride here.
Frontend	125	0.63	25	25	25	25	25	The frame around every screen, the patient overview, the assistant entry points. This is the role that paces the phase.
QA	85	0.43	15	15	15	20	20	A decision about how much verification the pilot buys, not a disagreement about the work.
Project manager	75	0.38	15	15	15	15	15	One contact on our side for the whole window, as you asked for.
Platform and DevOps	32	0.16	20	12				The environment first, then quiet.
Architect · SA/TL	25	0.13	5	5	5	5	5	Five hours a week against the pacing role.

The team 522 2.61 120 112 100 105 85

Budget \$29,555 Fixed for the scope on this page

Five weeks, twenty-five working days, 2.6 full-time equivalents. Not six part-timers — a small team with one role paced against another. Rates and commercial terms come separately, so this stays a technical statement you can check against the scope it is quoted for.

The one thing sitting outside this window: four to six days if consultation notes are confirmed for the pilot rather than left as a control that does nothing.

Clinexus V1.0 — five tranches, bought one at a time

Quoting V1.0 as a single figure would be useless to you: it is a number you cannot act on until the whole round has closed. Cut this way, each part delivers something that works on its own, and any one of them can be started, deferred or dropped without unpicking the others.

V1.0

TRANCHE	WHAT IT BUYS	PERSON-DAYS
Product	The application itself — the screens, the assistant, notes, oversight, localisation and the administration surface	136–217
Integrations	The adapters to your source systems, clinic identity, and the discovery that has to run against a real one	26–44
Finland AI runtime	The self-hosted model and speech recognition, and the Finland-resident environment they run in	46–73
Discovery	Answers bought before the work they gate. Decline any of them and the assumption it would have replaced stays yours	7–8
Regulatory	Our contribution into the technical file you own — hazard analysis, the evaluation gate, the software inventory	29–47

Clinexus V1.0 243–388

Budget Orientation figure — moves with the open questions \$117,950

The order that makes sense on the evidence we have: Product first, because it is what your clinicians and your investors see; Integrations in parallel with the conversations you are already having about source access, since those run longer than the build; the Finland runtime before any identified patient data exists; Regulatory alongside all of it rather than at the end, because a technical file assembled afterwards is assembled twice.

Three of the five parts depend on answers nobody has yet. That is why this is a range and not a price.

Why one is fixed and the other is not

The pilot can be taken at a fixed figure. It is specified tightly enough — the scope on this page draws the boundary explicitly — that we are comfortable committing, with an agreed mechanism for changes.

V1.0 should not be. Fixing a price against unknowns means pricing the unknowns, and you would be paying for our uncertainty. Time and materials by tranche costs you less and gives you a stopping point after each one. The range narrows as the open questions close.

Effort is in person-days and hours; rates and commercial terms come separately, so these pages stay a technical statement you can check.

What we need answered

How to read this section

These four sheets are the reason the V1.0 figure is a range. Almost every open item belongs to V1.0 rather than to the pilot — which is why we would rather start the pilot and close these alongside it than hold everything until they are closed.

Nothing here needs a meeting. Answering in writing, sheet by sheet, is fine.

Questions for Clinexus — where two of your statements cannot both hold

Contradictions, not doubts. Each card quotes both sides — your written answer and your own design — and puts our proposed reading underneath, so you can correct it in one line.

NEEDS AN ANSWER

Things we need you to settle

Every item below is a pair of statements that cannot both hold. Each quote is yours — from a document you sent or from the annotations on your own designs. None of them is a criticism: a product described in July and drawn in July will disagree with itself in places, and finding those places is what this reading was for. What we need is a decision on each, because seven of the nine change what gets built.

WHERE TWO OF YOUR OWN DOCUMENTS DISAGREE

No design is involved and no reading of ours is involved. Both quotes are written statements you sent us, and they cannot both be planned for.

● Role-based access control, for exactly one role

3-5 days if more roles are needed later

PRODUCT OVERVIEW v1.0 — 10 JULY

"Role-based access control"
Clinexus — Product Overview v1.0

ANSWERS Q17, ROLES AND TENANCY — 28 JULY

"Only the general physician role in pilot and V1.0"
Answers Q17

A role system that has exactly one role is a login. Both are perfectly buildable, and they cost differently — a permission model added later is more expensive than one designed in, but paying for it now buys nothing until a second role exists.

WE PROPOSE Build access control for the single general-physician role, and defer the role model to the phase that introduces a second role. If nurses or administrators appear in V1.0, that is a change we would want to hear about before January.

● A frontend-focused pilot, priced as a backend built from scratch

sets what the August date buys

CALL SPECIFICATION — 10 JULY

"August 2026 pilot, frontend and usability focused, backend not intense"
Clinexus call specification v1

RELAYED IN CHAT — 29 JULY

"Estimate from the Figma, costing all backend from scratch"
Message relayed through Denis

The first describes the pilot as a frontend exercise. The second asks us to price a complete backend. Both were said about the same three weeks in August, and the difference between them is most of the pilot budget.

WE PROPOSE A pilot with a working backend is not a frontend-only build. We have priced a real backend, because the screens cannot be demonstrated to clinicians without one — but we need you to confirm that this is what August is meant to deliver.

WHERE A WRITTEN STATEMENT DISAGREES WITH YOUR OWN DESIGN — ordered by what the answer changes

The right-hand quote in each of these is an annotation on your Figma file, and the screens named underneath are the artboards it appears on. Where our reading of a screen is part of the argument, we have said which screen, so you can check it.

● Read-only access, and screens that save

24-40 days, and the device classification

PRODUCT OVERVIEW v1.0 — 10 JULY

"Operate alongside existing EHRs with read-only access"
Product Overview v1.0

"Will not modify or write back to external clinical systems"
Product Overview v1.0

"Not intended to replace the EHR or perform clinical documentation"
Product Overview v1.0

YOUR FIGMA ANNOTATIONS — 28 JULY

"Generate clinical documentation on Stop Recording; review SOAP and approve actions"
Clinexus_V1.0 canvas

"Renew, edit, delete current medication; approve, edit, delete pending medication"
Clinexus_V1.0 canvas

On these screens: Current Medication table - New Pending Medication table - SOAP document - Send gate

See screens below: 1 2 3

Approving a medication, sending a referral and saving a note are all writes. If the product never writes to your clinical systems, those screens have nowhere to save to and the actions on them do nothing. If it does write, the product is performing clinical documentation — which is the one thing the overview says it does not do, and it is also the sentence a regulator will read first.

WE PROPOSE Read-only holds for the pilot, exactly as written. Every write path in the V1.0 designs becomes a separate, separately priced scope item, and we price it only once you confirm which of your systems will accept a write. This is the single largest open number in the estimate.

● Nothing generated beyond the record, and a screen that suggests what to do next

12-20 days, and whether the live feature exists

THREE DOCUMENTS — 10 AND 28 JULY

"Responses limited to the available patient record"
Product Overview v1.0, 10 July

"No free generation; nothing generated beyond what the system sees"
Call specification v1, 10 July

"Pilot Ask-AI is record lookup and summarisation only"
Answers Q2, 28 July

YOUR FIGMA ANNOTATIONS — 28 JULY

"Record live transcription; AI highlights data in real time and suggests future actions"
Clinexus_V1.0 canvas

On these screens: AI Suggestions - Key Findings - Case analysis answer

See screens below [2](#) [4](#) [6](#)

Finding something in the record and proposing what to do next are different acts. A suggested action is not in the record — it is produced, and producing it is what the three written statements rule out. The rule and the screen describe two different products, and only one of them is a Class IIa device with a simple story.

WE PROPOSE Tell us which rule governs V1.0. We have priced reasoning over the record with citations and abstention, because that is what the designs show. If the strict rule is the one you want, the suggestion features come out and this line roughly halves — but we would rather cut it deliberately than discover it at acceptance.

● Dictation only, and a transcript that separates two speakers

TI-17 days, and it decides whether anything leaves Finland

ANSWERS Q12, MICROPHONE — 28 JULY

YOUR FIGMA ANNOTATIONS — 28 JULY

"Microphone is dictation only; no audio storage; medical-domain ASR"

Answers Q12

"Record live transcription; AI highlights data in real time and suggests future actions"

Clinexus, V1.0 canvas

On these screens: [Diarised turn list](#) · [Recording transcript](#) · [Listening indicator](#)

See screen below [4](#)

A turn list separates the doctor from the patient, which means the whole consultation is being captured rather than notes being dictated. That is a different product to buy, a different vendor to choose, and it means audio of a patient exists — at least in transit — which is the only thing in the entire system that leaves Finland.

WE PROPOSE Dictation only in the pilot, on a bounded provider path. Two-party capture with speaker separation is a V1.0 item and priced there. We also need to know whether your Finland requirement covers speech as well as the language model, because that decision changes which vendors are even available.

● Alerts are read and never computed, and a panel that produces findings

blocks the alert integration

ANSWERS Q1, APPOINTMENT ALERTS — 28 JULY

YOUR FIGMA ANNOTATIONS — 28 JULY

"Appointment alerts read from the EHR; Clinexus computes nothing"

Answers Q1

"AI highlights data in real time and suggests future actions"

Clinexus, V1.0 canvas

On these screens: [Clinical Alerts](#) · [Key Findings](#)

See screens below [4](#) [5](#)

If nothing is computed, every alert on screen must already exist in the source system. The Key Findings panel shows findings produced while the consultation runs, and those cannot have been read from anywhere. Separately, we still do not know which resource in your EHR actually carries an alert — and without that, the tile has no production data behind it.

WE PROPOSE Alerts stay read-only as you wrote. We need one thing at kickoff: which resource carries an alert on the EHR side. Until we know, the alert path is priced against a single assumed resource, and findings produced during a consultation are treated as a separate feature, not as alerts.

● Referrals read-only in the pilot, and a design that sends them

the send destination is undefined

ANSWERS Q11, REFERRALS — 28 JULY

YOUR FIGMA ANNOTATIONS — 28 JULY

"Referrals read-only in the pilot"

Answers Q11

"Generate a referral letter, review, send, and see checklist gaps before sending"

Clinexus, V1.0 canvas

On these screens: [Send gate](#) · [Specialist selector](#) · [Letter draft](#)

See screen below [3](#)

The pilot displays referrals; the design writes and sends them. Sending needs a destination, and no destination has been named anywhere — not a system, not an address, not a protocol. It is the smallest of the write questions but it has the largest hole in the middle.

WE PROPOSE Referral display only in the pilot. Generation, the completeness checklist and transmission are V1.0 and we have priced them, on the assumption that you will name a destination before that work starts.

● Every fact should carry its source, and one flow does

the grounding claim depends on it

THREE STATEMENTS — 10 AND 28 JULY

YOUR FIGMA ANNOTATIONS — 28 JULY

"Displayed information should reference its originating source record"

Product Overview v1.0, 10 July

"Source and date shown on report cards"

Clinexus, V1.0 canvas

"Source references for Ask-AI answers; not planned for the overview"

Call specification v1, 10 July

On these screens: [Reports card in drawer](#) · [Answer body](#)

"Preserve source and date for every fact"

Answers Q10, 28 July

See screens below [5](#) [6](#)

Three statements set three different scopes: everything on screen, only the assistant, and every fact. The design carries provenance on one flow. This matters more than it looks — an assistant that cites its sources is a defensible product, and one that cannot show where an answer came from is not.

WE PROPOSE Source and date on every fact, as the most recent answer says. That also means adding two things the chat surface does not have: a citation affordance on each claim, and a way for the assistant to say it does not know. Both are priced.

● **Three assistant designs — one product, or three?**

the interface is priced once

PRODUCT OVERVIEW v1.0 — 10 JULY

"A controlled interface for patient-specific questions"
Product Overview v1.0 — singular, one interface

Nothing written anywhere describes more than one assistant. The largest of the three flows — pre-screening, eleven screens — has no written requirement at all.

See screens below 5 6 7 8

Read side by side, these are not one component in three states. They share a text box and a collapsed rail; the headers differ, the opening states differ, and — the expensive part — an answer is rendered three different ways. In pre-screening a reply about medication comes back as a structured card; in AI Chat the same kind of reply is prose. That is a different feature, not a different style, and we cannot tell from the file which one you intend.

WE PROPOSE We have priced one assistant surface with three modes, assuming the composer and the rail are shared and the rest are states of the same component. If that is wrong — if these are three surfaces with three answer renders — the interface work is materially larger than the estimate carries. Confirm the reading, and tell us specifically whether an answer about medication should come back as a structured card or as text.

YOUR FIGMA — THREE SEPARATE FLOWS

AI Chat — 7 screens

greeting, starter prompts, expanded sidebar - answers render as text

AI Chat pre-screening — 11 screens

own header, opening brief, chip row - answers render as medication, lab and note cards

Rare Case Analysis — 8 screens

collapsed rail only - answers render in their own block

Shared between all three: a text box and a collapsed rail. Nothing else.

● **The overview reads clinical meaning on two thirds of the screens**

four features nobody has costed

THREE STATEMENTS — 10 AND 28 JULY

"Appointment alerts read from the EHR; Clinexus computes nothing"
Answers Q1, 28 July

"No free generation; nothing generated beyond what the system sees"
Call specification v1, 10 July

"Responses limited to the available patient record"
Product Overview v1.0, 10 July

See screens below 1 2

Each of these four is a reading of clinical data rather than a display of it. Deciding that a value is severe, that two medicines interact, that a trend is worsening or that a result is out of range is a judgement somebody makes. The full argument, with the law behind it and the two products the answer picks between, is on its own sheet — Read or derive.

THE CHEAPEST WAY IN — three questions about one panel

Latest Summary — see screen 2

That panel groups findings into Key Complaints and Neurological Findings, gives every chip an information affordance, and highlights one chip in each group in a different colour. We cannot tell from the file what any of that means, and three short answers would settle the whole item above.

What does the colour mean, and how many values does it have?

We read blue and green as the two groups and red and orange as a highlight — but if green means normal, then weak reflexes and tremors in green is odd. We cannot tell whether there are two states or one.

What is behind the information icon on each chip?

The record entry the chip came from, or an explanation of it? Provenance is a link. An explanation is generated text, and that is a different feature under a different rule.

Who decides the two groups?

A fixed pair of sections, or does the product choose the categories for each patient? A dynamic taxonomy is a much larger thing than a two-section layout.

WE PROPOSE Tell us where these four come from. If your source systems already carry the severity level, the interaction flag and the reference range, then we integrate and display them, and we price that integration. If the product is meant to work them out, it is interpreting clinical data — which changes the estimate, and changes the regulatory conversation rather more than it changes the estimate.

● **Screens the product needs that your design does not contain**

we design as well as build them

CALL SPECIFICATION — 10 JULY

"No UI/UX design work from Light IT"
Clinexus call specification v1

We have taken this at face value and priced no design work. These seventeen screens are the exception it did not anticipate.

Your designs cover what a clinician touches, and they cover it well. What they do not cover is everything around it: signing in, what the screen does when a source is slow or wrong, and the operator screens somebody needs in order to run

SCREENS THAT MUST EXIST AND DO NOT

Four in the pilot

Sign-in - session expiry - degraded-source states - error and recovery

Six for operators in V1.0

audit log viewer - section-to-source configuration - model and prompt versions - retention and deletion - clinic users - language selection

Five more in V1.0

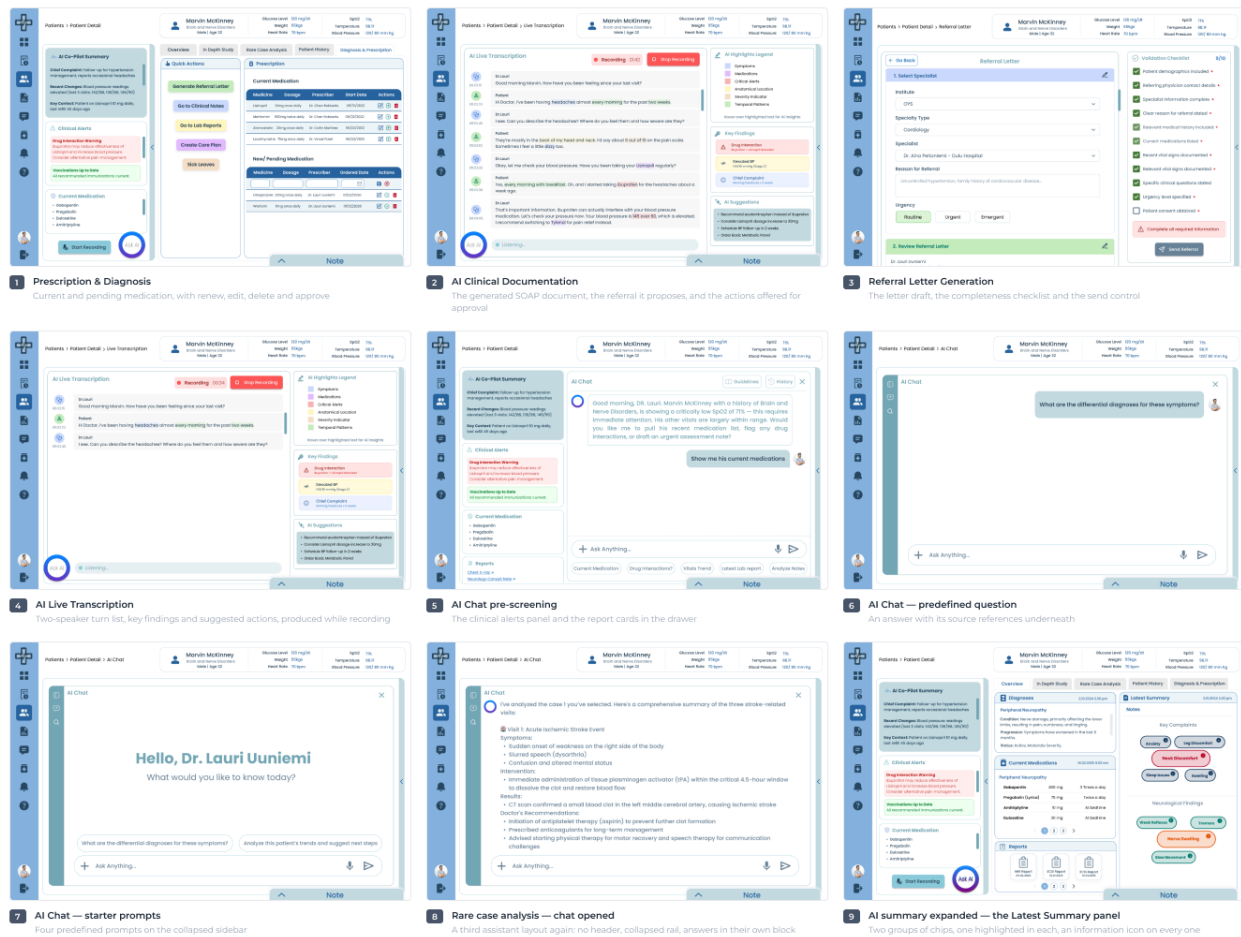
consent capture - export and print - document viewer - consultation note editor - clinical reference browser

the product — see the audit log, change which system feeds which section, roll back a model version, action a deletion request. Six of the seventeen will never be seen by a doctor at all.

WE PROPOSE We design and build these seventeen ourselves and have priced them that way, which is the one place your no-design-work instruction cannot hold. If your designer would rather draw the six operator screens, tell us before V1.0 starts and we will take that work out.

THE SCREENS THESE QUESTIONS COME FROM — your own artboards, unedited

Taken from Clinexus_V1.0 as delivered on 28 July. Nothing has been edited. Each card above names the screens it refers to by the numbers below. Screens 5 to 8 are the four assistant layouts, side-by-side.



What we need, and in what order

Two answers unblock most of the rest. Whether the product may write to your clinical systems decides four of these nine items and the largest number in the estimate. Whether V1.0 may reason over the record, or only look things up, decides two more and the shape of the live-consultation feature. Everything else can wait for kickoff.

None of the nine is a blocker to starting. They are blockers to being right about what August and January contain, which is a different and cheaper problem to solve now than in December.

Every quotation is verbatim from a document you sent or from an annotation on your own Figma file. Dates are the dates we received them.

Light IT Global

The sharpest one: a tab offering medication add, edit and delete, against your statement that the application writes nothing back.

Where the design outruns the description

Screens that exist in the designs and that no document explains. These carry our assumptions rather than proposals — we had to assume something to price it, and here is what we assumed.

NEEDS AN ANSWER

Where the design outruns the description

A different kind of question from the contradictions sheet. Nothing here disagrees with anything — these are places where you have drawn a working screen and no document anywhere says how it works. We can build to an assumption in each case, and have, but the assumptions are load-bearing and one of them is the largest regulatory decision in the product.

● Rare Case Analysis — eight screens built on one sentence

four different systems fit this description

THE ONLY THING WRITTEN ABOUT IT

"View rare cases of other patients and analyse similarity with the current patient"
User story on your Figma canvas — the whole specification

There is nothing else. No answer, no call note, no line in the overview describes how similarity is found or what the output is meant to be.

See screen below **1**

Retrieval over a curated corpus, structured matching on diagnosis codes, and a model asked to reason about similar patients are three different systems — different cost, different infrastructure, different regulatory position. The screen showing a prognosis is a fourth thing again, because a prognosis is not in any record: it is produced. And age-and-sex is not de-identification here, since the whole point of a rare case is the clinical detail that makes it rare — which is also what makes the patient findable.

WHAT WE ASSUMED Retrieval over a corpus that you assemble, approve and own, with the output cited back to the cases it came from and no prognosis. Priced that way, and held outside the estimate until the lawful basis for a cross-patient corpus under Finnish secondary-use rules is established. If you intended something closer to the screen, we need to talk before it is built rather than after.

WHAT THE EIGHT SCREENS SHOW

A list of other patients
labeled as age and sex only

A tray the doctor picks cases into
a selection set for analysis

An Analyze action
runs a comparison over the picked set

A per-visit narrative
with recommendations and a prognosis

● The largest flow in the file has no written requirement at all

eleven screens taken on trust

WHAT IS WRITTEN ABOUT PRE-SCREENING

"A controlled interface for patient-specific questions"
Product Overview v1.0 — the same sentence that covers the whole assistant

Pre-screening is the longest single flow you have drawn. Nothing distinguishes it from the ordinary chat in any document.

See screen below **2**

We have read this as a guided intake before the consultation, and priced it as modes of the assistant. But it is the flow with the most screens and the least description, and the artefact cards in it are a rendering feature that appears nowhere else. If pre-screening is meant to be its own product surface rather than a mode, the interface estimate is short.

WHAT WE ASSUMED A guided mode of the same assistant, sharing its service and its grounding rules, with three artefact renderers added. Confirm that reading — and if pre-screening has a clinical protocol behind it that we have not seen, that protocol is the thing we most need.

WHAT PRE-SCREENING ADDS THAT CHAT DOES NOT

Its own header and opening brief
a guided start rather than a blank prompt

A chip row
a structured way to narrow the question

Medication, lab and note cards
answers rendered as clinical artefacts, not prose

● Three kinds of note, and the screen that would tell them apart was never drawn open

one of the three may be a write

WHAT IS WRITTEN ABOUT NOTES

"Notes easily accessible, saved per consultation"
User story on your Figma canvas

The Doctor Notes flow is a single frame, and the note dock is drawn collapsed on every screen it appears on.

See screens below **1 2**

These read as one feature and behave as three. The important part is where the middle one lands: if a note typed in our interface is meant to appear in the patient record, that is a write to your clinical system and it belongs with the read-only question rather than here. If it stays with us, then we are holding clinical text that your record does not have, and somebody has to decide how long we keep it.

WHAT WE ASSUMED The doctor's note is ours, stored against the consultation, and never written back. The SOAP note is a separate artefact from it. Neither is the same as the clinical notes we read from your record.

THREE DIFFERENT NOTES EXIST IN THE PRODUCT

Clinical notes from the record
yours, read-only, already in the EHR

The note the doctor types here
ours to store — or yours, if it goes back

The generated SOAP note
produced from the consultation, and stored

● The specialist selector needs a directory that nothing provides

an unscoped integration, or a list somebody maintains

WHAT IS WRITTEN ABOUT REFERRALS

"Generate a referral letter, review, send, and see checklist gaps before sending"
User story on your Figma canvas

Nothing anywhere mentions where the choices in the selector come from.

See screen below

3

A three-level picker of institutions, specialties and named clinicians is a directory. It has to be read from somewhere, and no source in the plan supplies it. We reserved the FHIR resources that would carry it during the estimate reconciliation and never mapped anything onto them — which is our omission, but the source is still your decision. The pre-filled reason is a smaller version of the same question, and the checklist header disagreeing with its own list is worth a look on your side.

WHAT WE ASSUMED A directory read from the same source system as the rest of the record, using the standard FHIR resources for practitioners and their roles. If Finland has a national or regional directory you expect us to read instead, that is an integration nobody has scoped.

WHAT THE REFERRAL SCREEN NEEDS

A cascading picker
institution, then specialty, then a named specialist

A suggested referral reason
free text arriving pre-filled — suggested by what?

A completeness checklist
eleven criteria, and the header on the screen reads 9/10

● Four lists on the screens are clinical decisions, and no document makes them

content, not configuration

WHAT THE SCREENS SHOW

Three urgency levels
Routine, Urgent, Emergent on the referral form

An eleven-point completeness check
run before a referral can be sent

Six vitals in the patient banner
styled to show when a value is out of range

Six highlight categories
colouring what the assistant finds in the transcript

Every one of these is a clinical judgement rendered as a control. We can build all four, but we should not be the ones deciding them — and once the product is a certified device, changing what reads as abnormal is a release rather than a setting.

WHAT WE PROPOSE You supply the four as clinical content and sign them off; we build and version them with the device. One is already concrete: the form offers Routine, Urgent and Emergent, while the FHIR standard carries routine, urgent, asap and stat — with no Emergent — so a referral sent back to your system needs an agreed mapping.

WHAT NO DOCUMENT SETTLES

Which levels exist and what each one obliges
and whether they are yours or the receiving units

Which eleven criteria, and who chose them
the header on the screen reads 9/10, so even the count is unsettled

Which six vitals, and what reads as abnormal
the ranges are a clinical call, not a display option

Which six categories, and where the boundaries fall
this one shapes what a clinician's eye goes to first

● Some of the words on the artboards are the product, and some are filler

two sentences from you settle both

WHAT IS PAINTED ON THE SCREENS

Thirteen past conversations in the expanded chat sidebar

Suggest ICD-10 · Care plan generation · Discharge instructions · Compare treatments · Drug interaction · Caps in preventive care · Lab trends analysis, and six more

Six vital values in the patient banner, on every screen
glucose, weight, heart rate, SpO2, temperature, blood pressure — see screens 1 and 4 below

We have read both as sample content. That is a reading, not a fact — and it is the cheapest thing on this sheet to correct, because one sentence each settles it before anyone builds to the wrong assumption.

WHAT WE ASSUMED The chat history is designed sample content and the vitals are sample values. If either was meant as specification, tell us which lines — they become scope with a price, rather than a misunderstanding found during the build.

WHY IT CHANGES THE ANSWER

If the thirteen are a backlog, thirteen capabilities are missing

If they are sample content, nothing is missing. They appear in no document anywhere else, only on that one panel

If a banner value is a real threshold, it is a clinical setting
if it is a sample, building to it would be a defect rather than a feature

● Pending medication is drawn, and the resource that should feed it maps to nothing

a mapping question, not a scope one

WHAT IS WRITTEN

"Renew, edit, delete current medication; approve, edit, delete pending medication"
User story on your Figma canvas

Current medication has a source resource and a mapping. Pending medication does not.

See screen below

4

This is the smallest item on either sheet and it is here because it is cheap to close. Current medication reads from one FHIR resource; pending medication is a different resource and we have not confirmed that yours carries it in the way the screen assumes.

WHAT WE ASSUMED Pending medication comes from MedicationRequest and current medication from MedicationStatement. Confirmed at the mapping workshop with the rest of the section-to-source rules.

THE GAP

A pending-medication table
drawn, with approve and edit actions on each row

MedicationRequest
the resource that should carry it — reserved, never mapped

● The pilot has three possible data sources and we priced one

decides what August proves

ANSWERS Q5, INTEGRATION PATH — 28 JULY

"LAMK synthetic profiles, public Epic sandbox, ESKO unconfirmed"
Answers Q5 — recorded by us as conditional, meaning not settled

You also told us there are no API documents, credentials, sandbox access or endpoints available yet. So none of the three is

THREE CANDIDATES, AND WHAT EACH ONE BUYS

Our own synthetic FHIR server

Craspest and fully under our control. Proves the workflow and the interface. Proves nothing at all about integration, because it always answers.

reachable today.

The public Epic sandbox

The only candidate that proves the FHIR path is real. Costs an external dependency inside a three-week build, and we have no access yet.

LAMK synthetic profiles

A Finnish dataset named in your answer that we have never seen. We cannot size it.

This is the one item on either sheet where we did not assume so much as decide. We stood up our own synthetic server, priced it, and drew it on the architecture as though it were settled — and your answer names three candidates and calls the integration path unconfirmed. The choice matters more than its cost: it decides whether August demonstrates a working product or a working integration, and those are different things to show an investor.

WHAT WE ASSUMED Our own synthetic server, on the reading that August is a workflow and usability pilot. If it is also meant to prove the integration, the Epic sandbox is the candidate and we would want that decided now rather than in week two — we have no access to any of the three today, and obtaining it is not our decision to make.

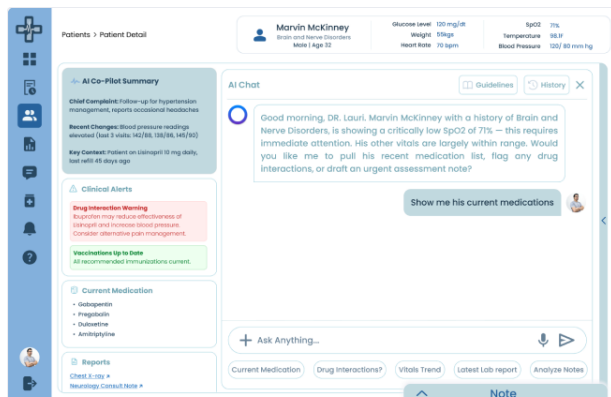
THE SCREENS THESE QUESTIONS COME FROM — your own artboards, unedited

From Clinexus_V1.0 as delivered on 28 July. Each card above names the screens it refers to by the numbers below.



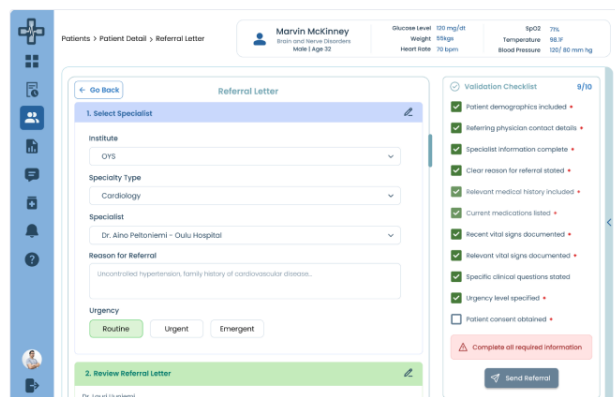
1 Rare case analysis

The list of other patients, the tray the doctor picks into, and the Analyze action



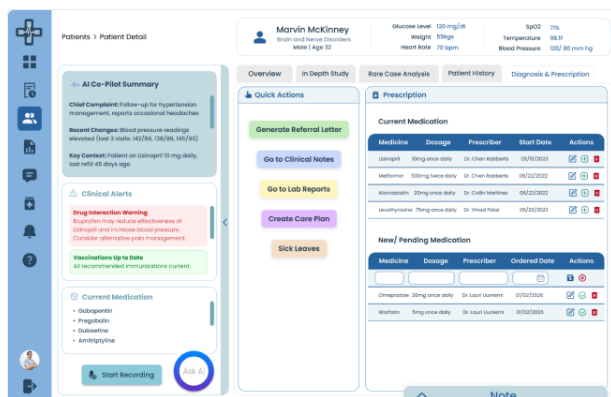
2 AI Chat pre-screening

Its own header and opening brief, the chip row, and answers rendered as clinical cards



3 Referral Letter Generation

The cascading specialist picker, the pre-filled reason and the completeness checklist



4 Prescription & Diagnosis

The pending medication table, with approve and edit on each row

Why these are on their own sheet

The other sheet asks you to choose between two things you have already said. This one asks you to say something you have not said yet. Every item here is a place where the drawing is ahead of the writing — which is normal in a product at this stage, and is exactly what a design file is for.

We have assumed an answer to each and built the estimate on it, so nothing here blocks a start. But an assumption is only cheap while it is still an assumption. Rare Case Analysis is the one that will not stay cheap.

Each quotation is verbatim from your Figma annotations or from the Product Overview. Where we say nothing is written, we mean nothing in any document you have sent.

Light IT Global

Rare Case Analysis is the largest of these — the screen exists and nothing describes how it works.

Screens that must exist and are not drawn

Sign-in and everything around authorisation, error states, the audit viewer, user administration, source selection, language, settings. Each traced to what requires it.

NEEDS AN ANSWER

Screens that must exist and are not drawn	
Your designs cover what a clinician touches, and they cover it thoroughly. What they do not cover is everything around it — signing in, what a screen does when a source is slow or wrong, and the screens somebody needs in order to run the product rather than use it. Almost every screen listed here is required by something you asked for in writing; each entry says which.	
"No UI/UX design work from Light IT" Clinexus call specification v1 — 10 July We took this at face value and priced no design work anywhere else. These screens are the one place it cannot hold: they have to exist for the product to function, and there is nothing to build from. Either we design them, or you do — but somebody has to, and it is cheaper to decide that now than in December.	
● Four of them are needed for the August pilot These are the urgent ones. The pilot is about five weeks of work and four of its screens do not exist in any form.	
Sign-in You asked for "secure user authentication" and "authentication compatible with healthcare environments". Every flow in your file starts with the doctor already signed in.	Session expiry and re-authentication You asked for "session management". A clinical session that times out mid-consultation has to do something, and nothing describes what.
Degraded-source states You asked that "missing, conflicting or outdated data" be "handled gracefully with a warning". Every screen in the file shows complete, clean data.	Error and recovery You asked for "reliable operation with appropriate error handling". There is no error state anywhere in the designs.
● Six of them no clinician will ever open Screens for running the product rather than using it. They are the reason an administration surface exists in the architecture at all — and it appears in none of the drawings, ours included, until now.	
Audit log viewer You asked us to log authentication, patient access, generated overviews and queries. A log nobody can read is not an audit trail — somebody has to be able to answer who saw what.	Section-to-source configuration You asked for "explicit agreed per-section FHIR resource and precedence rules". Once there is more than one source, the rules have to be visible and changeable without a release.
Clinic user administration You asked for role-based access control. Someone has to add and remove the clinicians who have it.	Model and prompt version administration A model or prompt change alters what the product says to a doctor. Under the device rules that is a release event with a record, not a setting — and it needs a place to be seen.
Retention and deletion administration A deletion request under GDPR has to be actioned by a person against a real record. Today there is nowhere to do it.	Language selection You asked for multilingual localisation. Something has to choose the language.
● Five more are clinician-facing and simply missing These sit inside flows you have drawn, at points where the flow needs a screen and stops.	
Consent capture and record Your referral send button is disabled until patient consent is given. Nothing anywhere captures or records that consent.	Consultation note editor The note dock appears on every screen and is drawn collapsed on all of them. Its open state has never been designed.
Document viewer The reports card in the drawer implies a document opens. Nothing shows what opening one looks like.	Export and print view Export, print and download are the points where clinical data leaves the product. They have controls and no destination screen.
Clinical reference browser You asked for approved external references "kept separate from patient data". Separate means a surface of their own.	
● Two are beyond V1.0 and listed for completeness Neither is priced. They are here so that the list is the whole list.	
Clinic onboarding and tenant configuration V1.0 is a single clinic. A second one needs a way to be set up.	Nurse vitals capture Your own stated ambition: nurses record vitals, written to the record on doctor approval. Not in V1.0, and worth naming as a direction.
What we have assumed, and what we would rather you decide We have priced all of these as ours to design and to build, which is the single exception to your no-design-work instruction and the only one we have taken. The four pilot screens are not negotiable on timing — without them the August build has no way in and no way to fail safely. If your designer would rather draw any of them, the operator screens are the natural place to start: they are the furthest from the clinical work you have been designing, and taking them back would remove the largest single piece of design effort from our side. Tell us before V1.0 starts and we will take it out of the number.	
Each screen here is traced to the requirement, capability or work package that needs it. None was added because it seemed like a good idea. Light IT Global	

Tell us which of these you do not want and we take them out — that is the cheapest correction on this page.

Read or derive

The four clinical readings your screens make, and the two different products the answer picks between. This is the single question with the widest effect on both the figure and the device conversation.

NEEDS AN ANSWER

Some things on your screens are conclusions, not facts

WHETHER YOUR SYSTEMS REACH THEM OR THIS PRODUCT DOES CHANGES WHAT IT IS, WHAT IT COSTS, AND WHO ANSWERS WHEN ONE IS WRONG

A patient's record holds facts: a glucose value, a list of medicines, the date of a visit. Your designs also show things that are in no record — a chip turned red because a value is critical, an arrow saying a trend is getting worse, a warning that two medicines should not be taken together. Somebody had to work those out.

Either your own systems already worked them out and this product simply shows them, or this product works them out itself. Nobody has told us which. There are four of them, they are all below, and the three screens they live on are at the foot of this sheet.

A — THE FOUR

Each one is already drawn, on screens you have seen. None of them is priced, because we do not yet know whether building it means fetching a value or reaching a judgement.

Severity

A chip or an alert carries a level — and an explanation of why.

WHERE YOU SEE IT

The summary panel at the top of the patient view, and the alerts beside it — screens 1 and 2 below

WHY IT IS A JUDGEMENT

Someone has to define the levels and decide which one a patient is at. The screen also shows why, and a reason is written by whoever reached it.

Drug interactions

A warning that two medicines should not be taken together.

WHERE YOU SEE IT

The clinical alerts, and the key findings in the co-pilot — screens 2 and 3 below

WHY IT IS A JUDGEMENT

This one is checked, not received — the design says so. A check is a clinical judgement about a specific patient's medicines.

Trends across visits

An arrow saying a value is rising or falling, and how often a prescription is refilled.

WHERE YOU SEE IT

The co-pilot summary and the day's list — screen 3 below

WHY IT IS A JUDGEMENT

A trend is a comparison across visits that no single record contains. Somebody computes it and decides what counts as a change worth showing.

Lab reference ranges

A result flagged as abnormal or critical.

WHERE YOU SEE IT

The lab report, which is a separate artefact and is not among the screens below

WHY IT IS A JUDGEMENT

Ranges differ by lab, by method, by patient age and sex. Flagging a value means holding a position on which range applies.

B — THE SAME SCREENS, TWO DIFFERENT PRODUCTS

European medical-device law draws its line about here. Software that shows a record is not a medical device. Software that produces information a clinician uses to decide on a diagnosis or a treatment is one — and lands in a class that brings an audit, a technical file and a notified body with it. Three of the four above sit on that line, and which side they fall depends entirely on the answer.

The hospital's systems already decide

The severity, the interaction check and the flagged ranges arrive with the record. The product fetches them and puts them on screen.

- 1 The work is connecting to those systems and reading the values — and once we can see them, the cost stops being a guess
- 1 The product stays a viewer, which is the lighter side of that legal line
- 1 The clinicians who set those rules stay accountable for them, which is also the cheaper answer for everybody

The product decides

Nothing upstream carries them, so this software works out what is severe, what interacts and what is out of range.

- 1 This is not in the estimate and not in what we agreed the product would do — it is new scope, not a detail
- 1 Three of the four then meet the definition of a medical device, with the file and the audit that follow
- 1 Each one needs a named clinician to approve what it decides, and evidence, kept over time, that it decides it correctly

C — THREE ANSWERS SETTLE ALL FOUR

Not a workshop. One panel on one screen — the Latest Summary block at the top of the patient overview, with the coloured chips. It is screen 1 at the foot of this sheet.

1

What does the colour on that panel mean, and how many colours are there?

That fixes how many levels exist everywhere else on the product — and tells us whether the level arrived with the data or was chosen here.

2

When a clinician taps the (i) beside a chip, what should appear?

If it is where the value came from, the product is showing. If it is why the value was given, the product is explaining — and that is the moment it starts making the judgement.

3

Who decides the two groups the chips are sorted into?

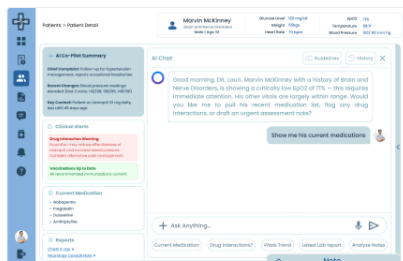
A fixed pair of sections, or does the product choose the categories for each patient? A taxonomy picked per patient is a much larger thing than a two-section layout.

THE SCREENS THIS IS ABOUT

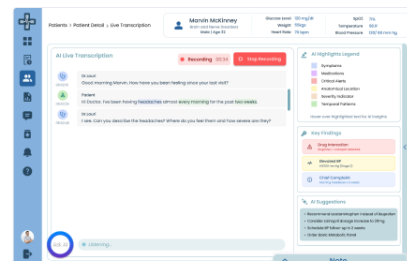
Your own airboards from Clinexus_V1.0 as delivered on 26 July, unedited. The lab report is a fourth artefact and is not shown here.



1 The Latest Summary panel
Chips with a colour and an information icon — the panel the two questions above are about



2 AI Chat pre-screening
The clinical alerts panel, with the drug-interaction warning, and the report cards in the drawer



3 AI Live Transcription
Key findings and suggested actions, produced while the consultation is still running

All four are drawn in the designs you sent us and none of them is in the estimate, because the same picture costs very different amounts depending on the answer. Nothing here is a criticism of the designs — it is the one question they cannot answer by themselves.

Light IT Global

Reading a value from the record and deriving a new one are not the same product, and not the same regulatory story.

