

WHITE PAPER

The Launch Decision Architecture

A framework for designing launches that hold up under year-one pressure. For global and regional HQs, and the country teams that run them.

By Pieter Vanderbeeken & Mark Watson, Co-Founders Fractal Force.

Key takeaways

The next three years bring one of the most crowded launch windows the industry has seen, with many of the largest companies bringing significant assets to market at once. They are competing for the same prescribers, in the same therapy areas, at the same time. In that commercial traffic jam, a launch built on share of voice (more reps, more messages than the competitor next door) hits a hard ceiling. **Relevance wins, not volume.**

When a launch misses, the molecule is rarely why. Of the 15 clear misses among the 50 most-anticipated launches of 2020 to 2025, only two failed on the science; the rest reached market with viable data and missed on the foundations underneath the launch: the data, segmentation, journey design, operating model and measurement the asset was supposed to run on. About two-thirds of launches miss their forecast, and in a launch there is no time to repair the foundations afterwards. **What happens in the first six months tends to stay where it lands.**

The Launch Decision Architecture is how we operationalise our Five Foundation Pillars under launch conditions: HQ owns **The Spine**, country teams pick **The Choice**, **The Cycle** moves both as the market does, and **The Glue** (a Launch Council and a Launch Cockpit) holds it together across markets and launches.

Get the foundations right and the year is winnable; get them wrong and no amount of spend will bend the trajectory back.

A snapshot of the launch wave that is forming now

Two facts about the next three years should concentrate the mind of anyone responsible for a launch. The industry is entering an unusually crowded launch window, and the organisations that will run those launches are in full execution mode or being assembled as this paper is read. A scan of the disclosed pipelines of the roughly 25 largest companies turns up 96 commercially significant launches across 2026, 2027, and 2028 in the US, EU, Japan, and China. The wave is large, and it is early: 61 of those 96 reach market in 2026, with the rest thinning out across 2027 and 2028.



Much of that density is a backlog clearing at once, and a striking share of these assets did not come from the launching company's own labs. They were acquired. Bristol Myers Squibb gained its schizophrenia drug, Cobenfy, through its acquisition of Karuna; AbbVie gained the Parkinson's asset tavapadon through its acquisition of Cerevel; Pfizer re-entered obesity by acquiring Metsera. For an established company moving into an adjacent area it already knows, that is routine, and the existing commercial base absorbs the new asset. The exposure comes when the asset lands in a genuinely new therapeutic area, where the company is building the medical relationships, the access approach and the field model at the same time as it launches. **Owning the molecule is not the same as owning the launch.**

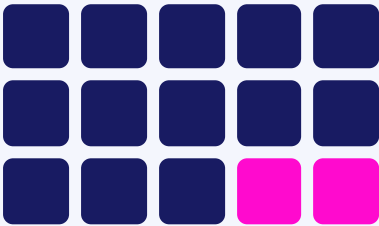
There is a second-order effect that matters more than any single launch. When 61 significant launches occur in a single calendar year, they compete for the attention of many of the same prescribers, across many of the same therapy areas, at the same time. Call it a "commercial traffic jam". In that environment, a launch built on share of voice (more reps, more messages, more frequency than the competitor next door) runs into a hard ceiling because the prescriber's available attention is finite, and every other launch is competing for the same slice of it. The launches that get through are the ones that earn relevance, not the ones that shout loudest.

What happened to the last wave?

What happens when a wave like this hits the market? We compared analyst consensus forecasts against reported performance for the 50 most anticipated launches of 2020 to 2025. A third of the launches old enough to judge missed outright, consistent with the long-run McKinsey research (see below). The hit rate is not the instructive part. What matters is why the misses missed.

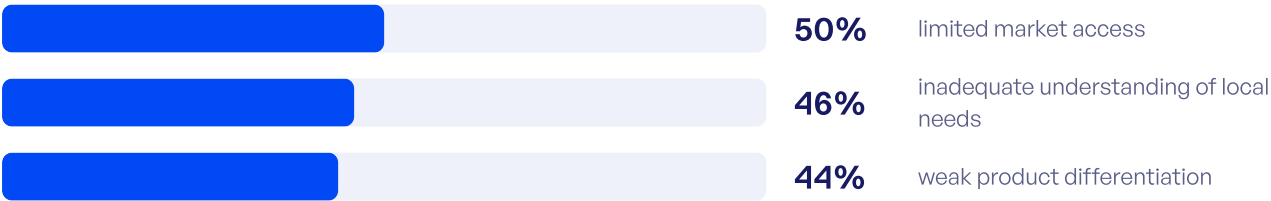
Of the fifteen clear misses, only two failed because the science did not deliver: Aduhelm, approved by the FDA despite its advisory committee voting against it and then effectively cut off when Medicare restricted coverage to clinical trials, and Relyvrio, withdrawn after its confirmatory trial failed. The other thirteen reached market with data that regulators accepted, and missed anyway. Where the gap is explained, the cause lies in the foundations: access and reimbursement that were not ready, an operating model that could not run the launch, a diagnostic pathway that did not yet exist, or a portfolio that cannibalised its own assets.

15 CLEAR MISSES – WHY THEY MISSED



- 2 failed on the science — Aduhelm, Relyvrio
- 13 reached market but missed on the foundations

Third-party data points in the same direction. Deloitte’s analysis of launches that missed their forecast found the recurring causes were limited market access (in 50% of misses), an inadequate understanding of local physician and patient needs (46%), and weak product differentiation (44%): three failures of foundation, not of molecule.¹ Trinity Life Sciences, working from a separate dataset, found that half of the 2023 US launch class underperformed its first-year forecast, an improvement on the 54% of 2020 to 2023, but still one launch in two.² Our own analysis of launches between 2020 and 2025 confirms the pattern: of the 41 launches in the cohort with a clear pre-launch peak forecast, 17 delivered less than 40% of it, and only 14 landed at 80% or above.



Two recent cases make the pattern concrete.

Sotyktu

BMS · PSORIASIS

\$291M

2025 sales, third full year — against a multi-billion ambition

Leqvio

NOVARTIS · CHOLESTEROL

<20,000

patients on NHS by 2025 — against an NHS target of 300,000

Bristol Myers Squibb's Sotyktu is a first-in-class psoriasis pill with strong trial data, yet it sold \$291 million in 2025, its third full year on market, against a multi-billion ambition.³ The explanations on record are foundation problems, not molecule problems: the company points to payer step-edits, the rules that force patients onto cheaper drugs first, while physician surveys point to lingering caution about the drug class's safety. Novartis's Leqvio is heading for a fraction of the \$2.5 billion analysts forecast for 2027; global sales reached just \$223 million in 2024.⁴ The brake was not the drug. Cardiologists in the US had to learn a new reimbursement and administration pathway the drug required, one they had no prior experience with, while waiting for outcomes data that would not read out for years. In the UK, NICE approved the drug and NHS England struck a first-of-its-kind national deal specifically to remove reimbursement friction — yet by 2025, fewer than 20,000 of an expected 300,000 patients were on treatment. The deal fixed one layer of the foundation. It did not fix patient identification and referral in primary care, which sat on a different layer entirely. Better foundations alone would not have closed gaps this size, and we do not claim they would; what the two cases show is where the ceiling came from. These are foundation gaps, measured in billions, on drugs as good as the ones the next wave is about to launch.

The pattern holds in reverse for the launches that beat their forecast. Wegovy and Mounjaro are the obvious names, and part of the honest answer is that both were so far ahead of anything else for weight loss that they would have sold heavily almost regardless. Even so, their early constraint was supply, not demand: Wegovy had to pause its own marketing in 2022 because Novo Nordisk could not make enough of it, and sales went vertical once capacity caught up. A genuinely differentiated drug can carry a flawed launch. Most launches are not that drug, which is exactly why the foundations decide the outcome for the others.

Launches are where foundations either hold or break

Most HCP engagement programmes underperform for a single underlying reason. The foundations they rest on were built for an earlier, brand-centric era of selling, in which the company decided what to communicate and the field force pushed that message through a small number of channels. HCPs do not care how a pharmaceutical company has organised itself internally. They expect to be engaged on their own terms, in their own clinical contexts, through the channels they choose, which now often include the AI assistants and large language models many use to answer clinical questions. The only thing that matters to them is the result: a conversation designed around them, or one that was not. The platforms (omnichannel suites, AI engines, content factories) layered on top of those brand-centric foundations cannot fix a base built on the wrong assumption: that the company, not the HCP, decides how the engagement runs. The result is the well-documented gap between digital investment and ROI that the industry has been writing about for a decade.

At Fractal Force, our work is to rebuild those foundations from brand-centric to customer-first, ensuring that the CX, omnichannel, AI, and other HCP engagement investments organisations have made finally start to deliver the ROI they were meant to deliver.

Foundations are vital in any HCP engagement setting. During a launch, they become even more critical. Not only is it challenging to fix the foundations while in full execution mode, but also, what happens in the first months on the market tends to stay where it lands. That is not just a marketing saying. The initial months are when HCPs develop their prescribing habits with the new asset, and research shows that doctors' adoption of new medicines remains consistent: prescribing habits strongly predict whether a new drug will be used at all, and adoption usually stabilises within roughly six months of launch. HCPs try the therapy on a few patients, observe how it performs for them, and form an attitude towards it—positive, negative, or indifferent. Once that attitude is established, the cost of changing it rises sharply. The clinical data are no longer novel. The conversation has taken place. The next time an HCP is genuinely open to changing their view on the asset is months away. The outcome data reflect the behaviour. In Deloitte's analysis of 284 US launches, of the drugs that did not meet expectations at launch, only a quarter reversed course in year two or year three, and fewer than one in five reversed it reliably across both years. A slow start is rarely a slow start that improves later; it is often the beginning of a trajectory the organisation will spend years trying to alter.

~6 months

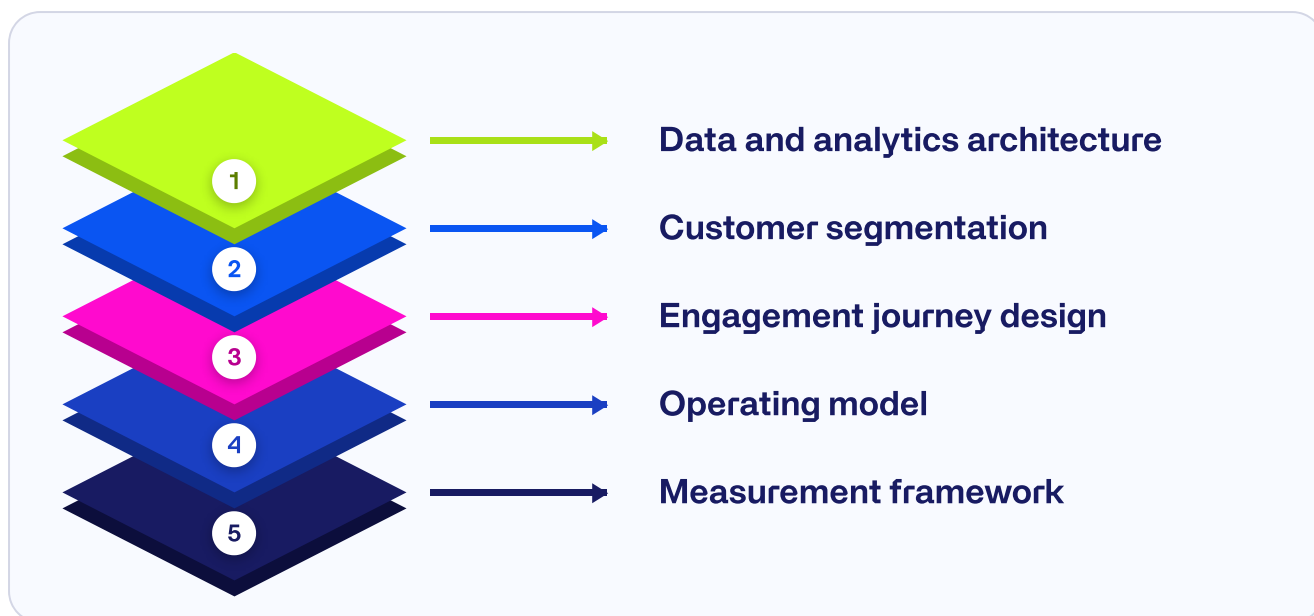
is all it takes for adoption to stabilise after launch —
the window in which prescribing habits form

1 in 4

launches that miss at launch reverse course in year
two or three; fewer than one in five do so reliably

None of this is new or contested. McKinsey's long-run launch research, the industry's reference point for a decade, found that about two-thirds of drug launches miss their pre-launch forecasts, a figure that has barely changed over the past 20 years.⁷ What the last wave showed, in the cases above, is that the misses came from the foundations the molecules were launched on, almost never from the molecules themselves.

What sits underneath every HCP engagement programme that works is the same set of foundations. We call them the **Five Foundation Pillars**: data and analytics architecture, customer segmentation, engagement journey design, the operating model that runs the engagement, and the measurement framework that says whether it actually works. In any HCP engagement setting, these are what determine whether engagement compounds across cycles or quietly disappears. In a launch setting, they determine whether the asset has a year-one trajectory that the organisation can live with.



No launch team needs to invent these Five Pillars from first principles. The work is deciding which parts of the Five Pillars HQ sets centrally and holds across every market, which parts country teams set locally for their specific conditions, and how both evolve as the market itself evolves. That is what we mean by the Launch Decision Architecture.

Three layers, held together by The Glue

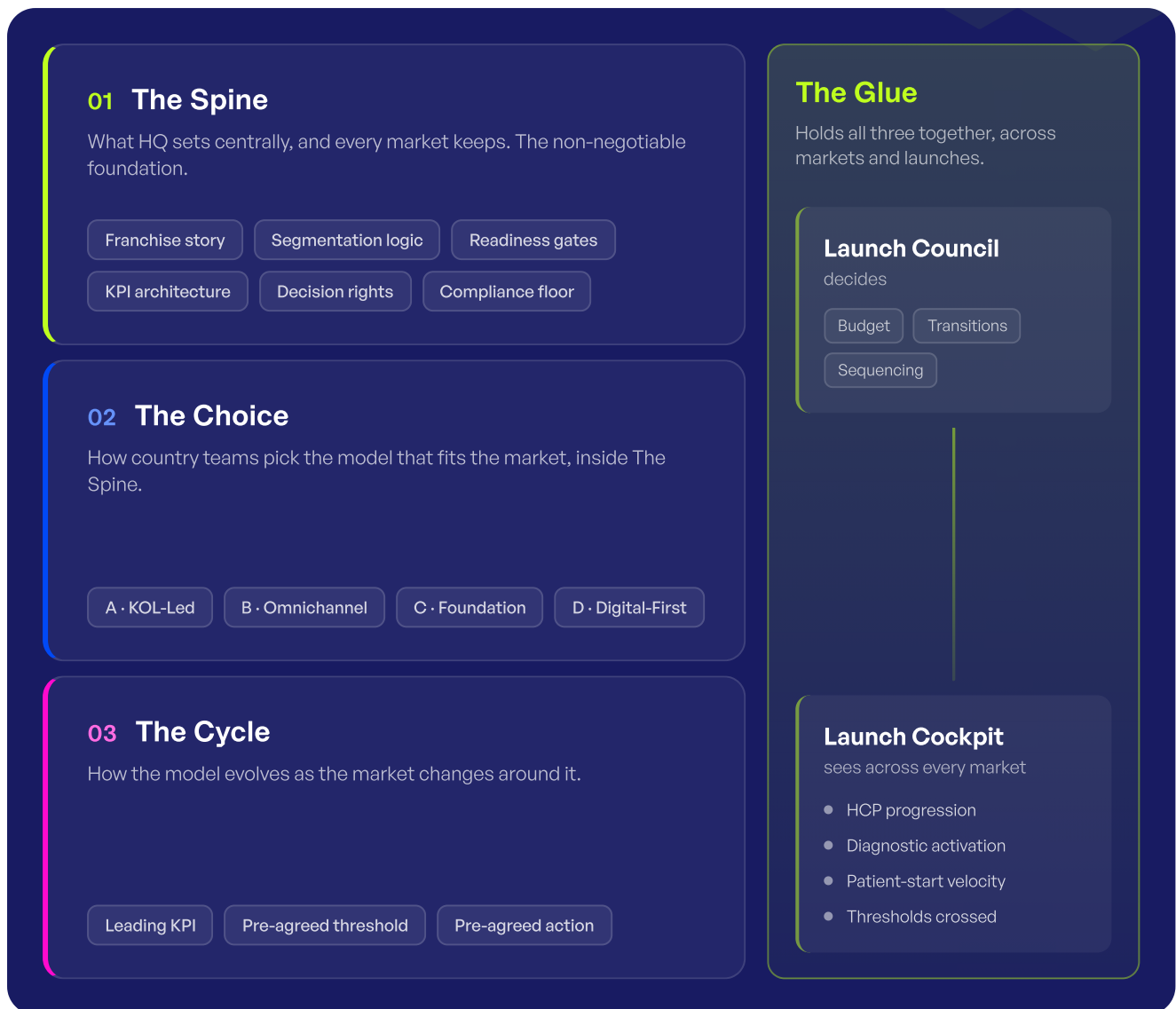
The Architecture has three layers, each answering a different strategic question, with a single governance system holding them together across markets and across time.

Layer 01 **The Spine** is what HQ sets centrally, and every market keeps without exception.

Layer 02 **The Choice** is the country team's choice to fit their specific market reality.

Layer 03 **The Cycle** is how the model evolves as the market changes around it.

The Glue Sitting across all three: a cross-functional Launch Council that decides, and a Launch Cockpit that lets The Council see what is happening across every market in real time.



Two clarifications before going layer by layer. The Architecture works the same way for a single brand launching into a single market and for a portfolio of brands launching across thirty. The variables change, but the structure remains the same. And the archetypes, scoring dimensions, and transition triggers in this paper are intended only as starting points. Real architectures are co-designed with the organisation that will run them, which is also the only way they get adopted on the ground rather than tolerated on paper.

How the Five Pillars sit inside the Architecture

Every pillar lives in a layer. The Spine fixes what HQ holds; the Choice and Cycle put the pillars to work in-market.

1 Data & Analytics Architecture THE SPINE + THE COCKPIT	One KPI architecture and the shared platform layer — the Cockpit reads it across every market.
2 Customer Segmentation SPINE SETS THE LOGIC · CHOICE APPLIES IT	A common segmentation method is a Spine constant. Local segments and tiers are a Choice-layer call.
3 Engagement Journey Design THE CHOICE + THE CYCLE	The archetype sets the journey at launch; the Cycle re-shapes it as HCPs progress through awareness and consideration.
4 Operating Model THE GLUE	The Launch Council and decision rights — who decides, and with what authority.
5 Impact Measurement THE SPINE + THE CYCLE	One measurement logic, set centrally; recalibrated to leading indicators when the model shifts.

The Architecture is how the Five Foundation Pillars get operationalised under launch conditions — each pillar mapped to the layer that owns it.

LAYER 01

The Spine

What HQ sets, and every market keeps.

Every Spine conversation hits the same first objection from country teams: “Our market is unique.” That objection is not wrong on the facts. Markets really do differ in terms of HCP access, healthcare system structure, regulatory posture, and partner footprint. What the objection misses is the point of The Spine. The Spine does not exist to crush local autonomy. It exists to define the sandbox in which country teams can act without having to ask HQ for permission at every step. Without that sandbox, local autonomy becomes a risk to the franchise. With it, local autonomy becomes the way HCP engagement gets stronger market by market.

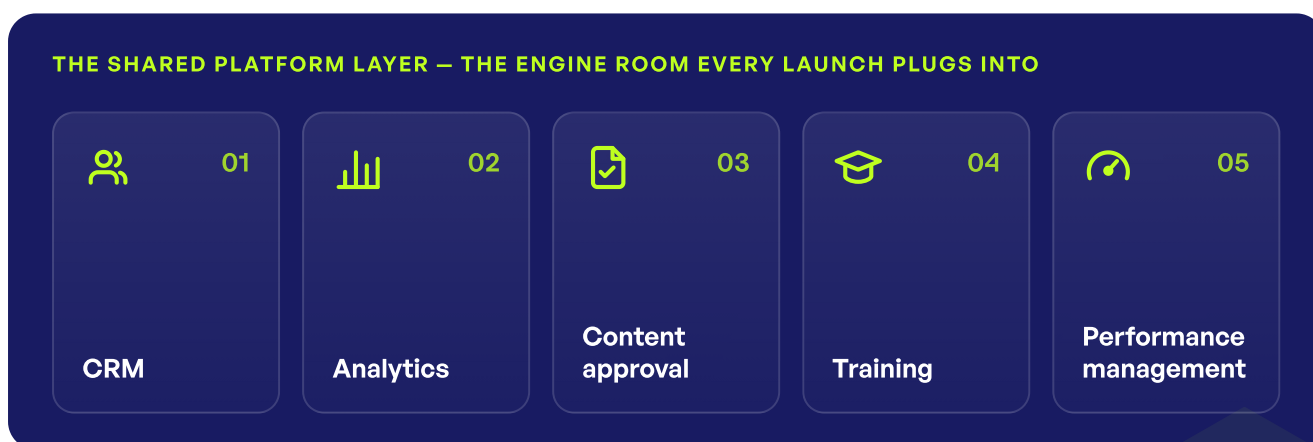
Concretely, The Spine is the non-negotiable foundation of the launch. HQ sets it once, and every market keeps it without exception. Its job is to protect global brand equity, ensure compliance, and improve operational efficiency for running multiple launches in parallel, while freeing country teams from reinventing the basics every time a new asset is launched. A weak Spine lets local autonomy degrade into rogue execution, where every market looks like a different company. A heavy Spine pushes country teams to stop bringing local intelligence back to HQ, because they assume the answer has already been decided. The discipline is to find the level of tightness that holds the franchise together without overpowering the local judgment that makes country teams useful in the first place.

Six constants do most of the work. They are the elements most consistently present in launches that hold together, and most consistently missing from those that fragment. Specific contexts can require more, but the principle does not change: what HQ sets is what every market keeps.

-
- 1 One franchise story.** Globally consistent clinical narrative, value proposition, and competitive positioning. For complex therapies such as antibody-drug conjugates (ADCs), gene therapies, and novel combinations, the safety protocols sit here as well. Local teams adapt how the story is told. The substance of the story does not vary by market.
 - 2 Common segmentation logic.** The same methodology for how HCPs and accounts get segmented and prioritised in every market. The actual segments and tier definitions can be local because the underlying HCP populations differ. The logic that produces them cannot deviate, because that is what allows learnings to move between markets and performance to be compared on the same basis.
 - 3 Unified readiness gates.** The same milestone checklist, the same readiness gates, and the same capability assessment in every market, regardless of whether the market runs full, hybrid, or lean. The Spine is where the organisation states clearly what readiness it will not launch without.
-

-
- 4 **One KPI architecture, feeding one Cockpit.** The same KPIs, the same measurement logic, and the same dashboards across every market and every model. One source of truth, with alerts that fire fast and feed The Council so that drift in one market can be acted on while it is still small.
 - 5 **Explicit decision rights.** Documented governance, escalation paths, and a clear decision-rights matrix between Global HQ, the affiliate, and any alliance or channel partner involved in the launch. A first-100-days checklist on every launch. Most Spine failures we see come from decision rights that were assumed rather than written down.
 - 6 **The compliance floor.** Medical information, adverse-event reporting, patient support, and market access capability. The quality and compliance floor that no efficiency choice may compromise, in any market, at any stage of the launch.
-

Underneath the six constants sits a shared platform layer: the CRM, analytics, content approval, training, and performance-management infrastructure that the launch actually runs on. The platform layer is the engine room that powers every launch the organisation will run, this year and for years to come. When the platform layer is in place, country teams plug into it the moment a new launch begins. When the platform layer is missing or fragmented, every launch rebuilds the same scaffolding from zero, and the operational learning that should have compounded across launches evaporates between them.



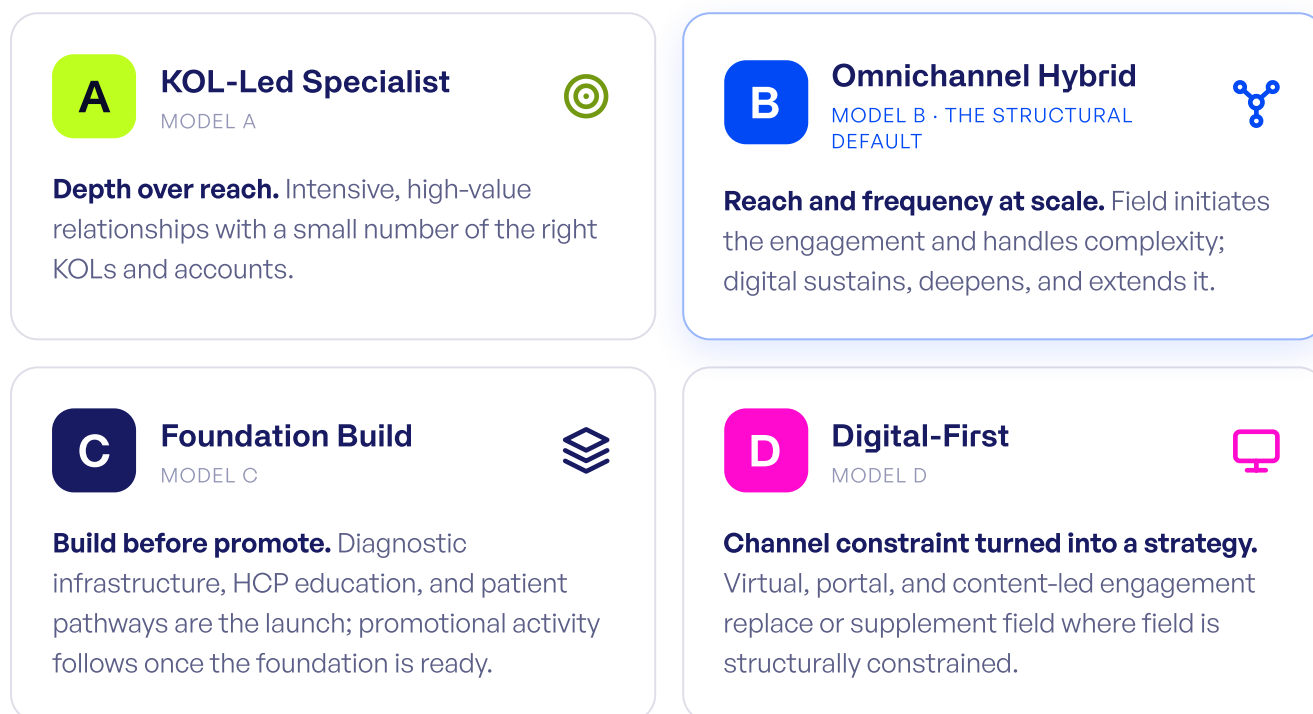
LAYER 02

The Choice

How country teams pick the model that fits the market.

The Choice layer is where the Architecture grants country teams real freedom to design their own launch for their market. The Spine establishes the parameters that every market must adhere to. Within those boundaries, the country team has full authority to select the HCP engagement model that best suits local conditions. HQ does not impose a model from the centre or dictate which archetype to implement. Instead, the framework offers a structured way to identify the most appropriate archetype for the market they are entering, along with a documented approach to justify that choice if the launch faces scrutiny later.

Four archetypes form the bounded menu of choices:



The diagnostic that gets a country team to the right archetype runs in two stages. Stage One uses four hard constraints to eliminate models that cannot work in this specific market, regardless of how attractive they look on paper. Stage Two uses eight scoring dimensions to identify the best fit among the remaining models. The full diagnostic (the four archetypes, both stages, and the scoring matrix) sits in the appendix at the end of this paper.

The diagnostic produces more than a model name. The real output is the documented rationale: the constraints, the dimension scores, and the judgment calls made when signals conflicted. That rationale matters in two specific situations that arise on every launch. The first is when budget owners or regional leadership push back on the model six months into launch, when early results are still developing. In that conversation, the rationale is what defends the original decision and turns it from an opinion into evidence. The second is when the launch team itself changes, as it usually does between launch and year three. In that handover, the rationale is what the next team inherits instead of starting the model conversation from scratch.

One principle is worth stating up front, because it shapes how many country teams approach The Choice layer. Omnichannel Hybrid (Model B) is the structural model for most established categories in most markets, and it remains the structural model as digital channels mature within it. Many country teams nevertheless treat Omnichannel Hybrid (Model B) as a transitional state on the way to KOL-Led Specialist (Model A), which they associate with a more committed, more prestigious launch posture. Digital, in that view, is the cheap backup. That instinct is one of the most consistent and most expensive misallocations we see.

“

Defaulting upward to KOL-Led Specialist because it feels more premium is the **single most expensive ego-driven misallocation in launch planning.**

The right answer in most established categories is to commit to Omnichannel Hybrid (Model B) as the deliberate, sustained choice it actually is, and to stop measuring launch ambition in field-force headcount.

LAYER 03

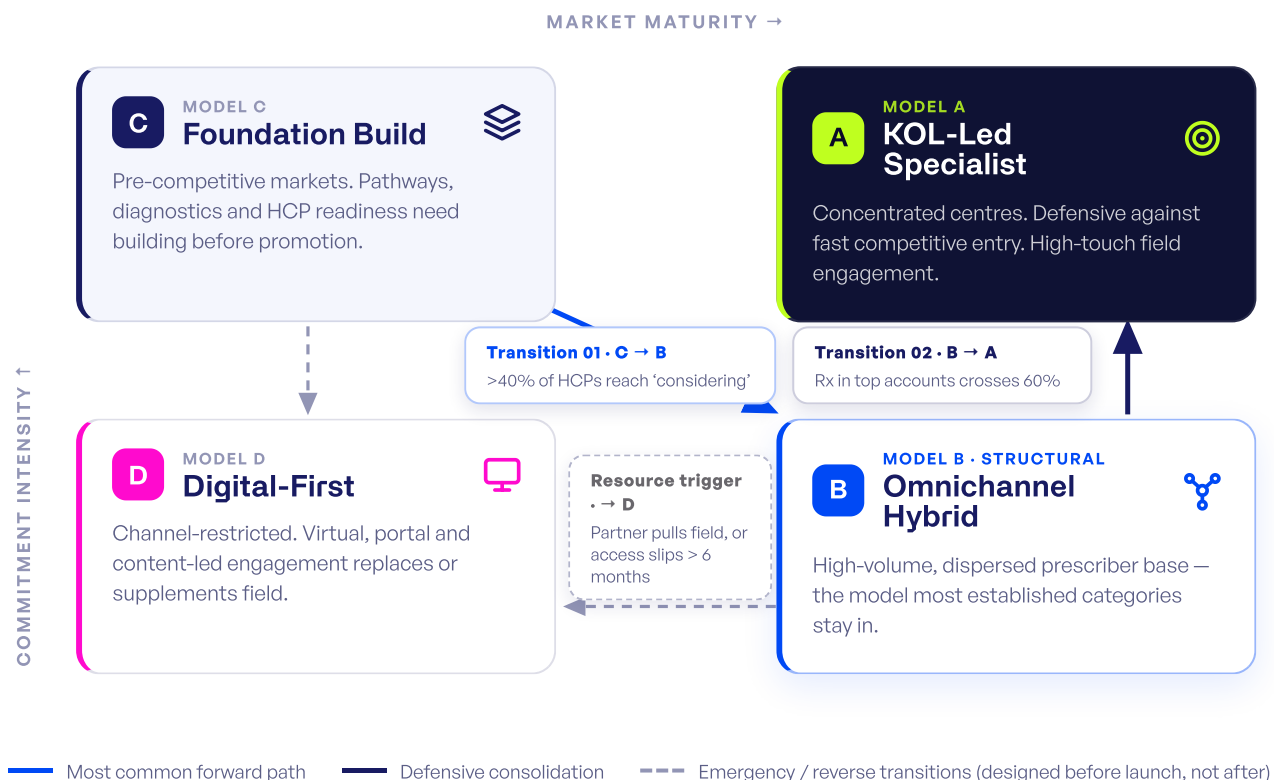
The Cycle

How the model evolves as the market does.

03

The model that fits at launch will not still fit eighteen months in. Indications expand, competitors enter, reimbursement decisions land or slip, and partner footprints shift. A meaningful share of launches that disappoint do so because the model was not adjusted quickly enough as the market began to move. The early signals that the model was no longer the right one were missed in the first months, and those signals did not improve on their own. The Cycle is the layer that makes course-correction operational rather than political. It is also the layer that most launch architectures handle the worst.

Most markets follow a recognisable arc



In a pre-competitive market with unestablished pathways, the right starting model is Model C, Foundation Build. As HCPs progress through awareness and consideration, and as the diagnostic and referral infrastructure matures, the market typically transitions to Model B, Omnichannel Hybrid, which is where most launches should stay for the long run. In a subset of markets where competitor entry into a concentrated account base creates real risk, the right move is a defensive shift to Model A, KOL-Led Specialist, where depth in the top centres becomes the defensible moat against share erosion.

Not every transition moves up the commitment ladder. The most expensive transitions are the ones that move backwards or sideways under pressure. A partner pulls field. A reimbursement decision slips by six months. A competitor's superior data forces an emergency rebuild of the engagement model. These backward and sideways transitions are exactly what the Cycle is designed to handle in advance.

A transition is three components, all agreed before launch

A transition comprises three elements: a leading-indicator KPI, a pre-agreed threshold for that KPI, and a pre-agreed operational action the organisation takes when the threshold is crossed. All three have to be in place before launch. The Council cannot bolt them on after the market has already started to move.

- 1 A leading-indicator KPI.** Prescription volume is a lagging measure. By the time it tells the organisation that the market has moved, the window for adjustment is closing. HCP journey progression, biomarker testing penetration, site activation rates, and patient-start velocity are examples of leading measures. They tell the organisation that the market is about to move, which is the window of opportunity for adjustment.
- 2 A pre-agreed threshold.** Numerical, defensible, and agreed before launch by The Council, the affiliate, and any partner involved. The benefit of agreeing the threshold in advance is psychological as much as mechanical. When bad news hits the launch, an organisation without a pre-agreed threshold tends to panic, debate, or default back to the playbook it ran before. An organisation with one executes the pre-agreed action. When good news hits, an organisation without a threshold tends to over-celebrate or hesitate to commit to the acceleration the data supports. An organisation with one moves faster because the next move has already been decided. The threshold is the ultimate governance tool: it replaces emotional alignment under duress with automated execution based on data.
- 3 A pre-agreed operational action.** An explicit, named change to the model: which content gets activated, which field roles get rebalanced, which channels open or close, and which capabilities the affiliate has to have ready in advance for the transition to be executable on the timeline The Council expects.

In the table, you can find three illustrative transitions that we see most often in practice. Real triggers are co-designed with the markets that will run them, in line with the realities of those specific markets. The structure is the same in every case: a leading KPI, an agreed threshold, and a costed, planned action.

#	LEADING-INDICATOR KPI AND THRESHOLD	MODEL SHIFT	PRE-AGREED OPERATIONAL ACTION
01	HCP journey progression. More than 40% of target HCPs reach the 'considering' stage on the engagement journey.	C → B	Activate digital channels. Shift field content from pathway education to clinical evidence. Add omnichannel reach to the existing account-based effort.
02	First prescriptions in target accounts. Adoption has crossed in more than 60% of them.	B → A	Deploy peer-to-peer programme. Shift field role to access support and prescribing confidence. Concentrate field effort in the top-decile accounts.
R	Resource trigger. Partner reduces field allocation, or reimbursement is delayed by more than six months.	Any → D	Pivot to virtual engagement. Launch HCP portal. Establish a regular webinar cadence. Reallocate the field-cost reserve to digital content and remote MSL coverage.

Where transitions actually fail: operational gaps

When a transition fails to land, the cause is almost always operational rather than strategic. The Council can formally approve the model change, the data can support it, and the leadership can communicate it across the affiliates. The problem shows up further down. The field force was trained for the previous model, and now needs new playbooks. The content library was built for the previous stage of the customer journey, and now needs to be reassembled for the new one. The approval queue is full of materials that the new model has just made obsolete. Each of those gaps is solvable on its own. Together, they typically add up to a quarter's worth of operational catch-up between the moment The Council changes the model and the moment the market actually sees a different launch on the ground. The Architecture has done its job at that point. The execution layer underneath the Architecture has not been prepared for what the Architecture has decided.

≈ 1 quarter

of operational catch-up typically sits between the moment The Council changes the model and the moment the market actually sees a different launch on the ground — unless the execution layer was prepared in advance.

Three operational disciplines must be in place before the transition is triggered, not assembled at the moment it fires.

- 1 Modular content, built once.** Build content libraries that can be recombined by stage of the HCP engagement journey rather than rebuilt at every transition. The same evidence modules, safety language, and access materials get assembled differently for the education, evidence, advocacy, and prescribing-confidence stages. The approval process approves the modules. Final assembly happens locally, which is what lets a market change its content mix in days rather than months when a transition fires.
- 2 A field operating rhythm designed to flex.** Train reps and MSLs for the role the next model will need, not only the role the current model uses today. An affiliate that has built deep account relationships under Foundation Build does not become an omnichannel team the moment the trigger fires for the shift to Model B. The new field roles need to be designed and partially rehearsed in the months leading up to the transition, so that the change is executable within weeks of the trigger.
- 3 KPI recalibration that fires with the transition.** When the model shifts, the KPIs the field team is measured on have to shift in step with it. A field force still being measured on call frequency after the transition to Model B has been triggered will keep running the previous model's playbook, because that is what the measurement system rewards. The Cockpit must update what is measured the moment The Council approves the transition, and the affiliate compensation and review system must update accordingly.

The Glue

A Council that decides. A Cockpit that sees.

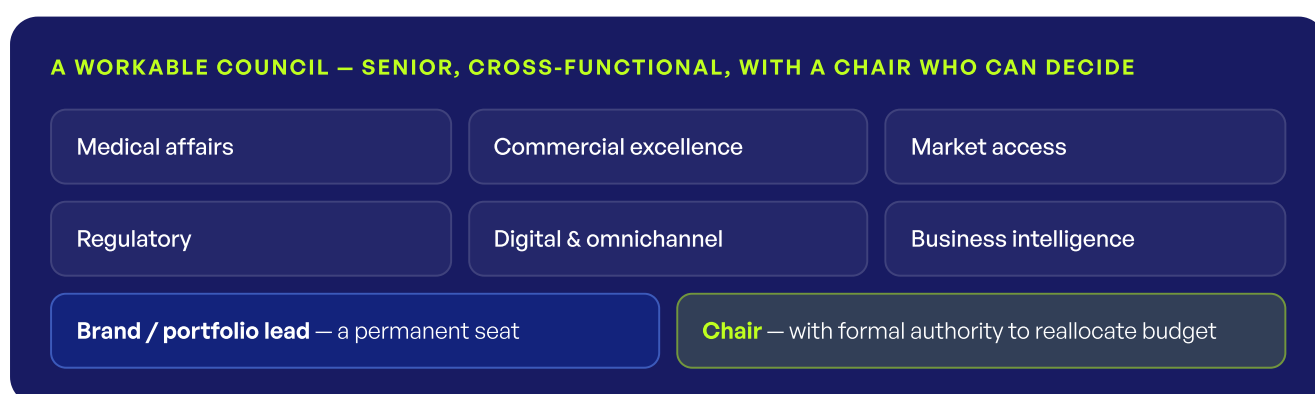


The three layers function only when a governance system holds them together across markets and across launches. Without that governance, The Spine fragments under affiliate pressure, The Choice drifts toward whichever model is politically easiest to defend in the moment, and The Cycle never actually fires because nobody has the standing to call the transition. The Glue does two things at once. The Launch Sequencing Council decides. The Launch Cockpit shows The Council what to decide on, and where in the launch portfolio attention is most needed right now.

The Launch Sequencing Council

The Council is a cross-functional body that co-owns launch decisions across medical, commercial, market access, regulatory, and digital functions, with explicit decision authority over the launch portfolio rather than over any single market in isolation. Its purpose is broader than approving individual transitions or arbitrating between markets that disagree. The Council is the body that sees the launch portfolio across all markets and all brands at once, and its job is to make every launch in that portfolio better than it would have been on its own. Co-developed launches make this concrete. When two organisations share an asset, as AstraZeneca and Daiichi Sankyo do across their antibody-drug-conjugate franchise, the launch runs on one shared set of decision rights that either holds or fractures under pressure. The Spine is where those rights get written down, and The Council is where they get exercised.

A workable Council has senior people from medical affairs, commercial excellence, market access, regulatory, digital and omnichannel, and business intelligence, a permanent seat for the brand or portfolio lead, and a chair with formal authority. The chair is the test of whether the organisation has built a Council or an advisory group. Real Councils have authority to reallocate budget across markets when the data justifies it, to redirect HQ-level resources to the market that most needs them, and to commit the organisation to model transitions on a defined timeline. Without budget reallocation authority, the framework still produces decisions; the organisation just executes them more slowly and at higher political cost.



Concretely, The Council does five things.

- 1 **Portfolio oversight.** It oversees the launch portfolio across all markets and all brands. It sees which markets are tracking ahead of plan, which are drifting, and where shared issues across markets need a coordinated response that no single market can produce alone.
- 2 **Support for local decision-making.** It supports country teams that are running workshops, executing transitions, or pushing through unusual market conditions. The Council is not a gate between local teams and their own market judgment. It is a resource the local teams can pull on when the decision they are about to make is one The Council has seen before.
- 3 **Cross-market learning.** It captures and circulates learnings from the early-launching markets so that the next-to-launch markets benefit from what the first ones discovered. This is the single largest source of compounding value across a launch portfolio, and the single most consistent place where organisations leave value on the table.
- 4 **Shared resource allocation.** It allocates the resources that no single market has access to on its own: specialist support from HQ functions, content templates that can be localised, capability investments, partner engagement at the global level. The Council places those resources where the portfolio needs them most, not where they are requested loudest.
- 5 **Decisive intervention when a market is drifting.** It acts decisively when a market is drifting away from plan, and the early signals show that the drift is not self-correcting. This is the role The Council is most often expected to play in name, but it is one of five. It works only when the four other functions of The Council are already in place.

Approving individual model transitions and resolving market-level conflicts sit within The Council's remit, but its centre is cross-market learning and portfolio-level oversight at a scale no single market can manage alone.

The Launch Cockpit

The Cockpit is the single source of truth across every launch and every market. One dashboard, one set of definitions, fast course-correction. The Cockpit tracks the leading indicators that matter for cross-market visibility:



Most launches that disappoint do so because the signal that something was off arrived later than the moment when something could still be done about it. IQVIA's Launch Excellence work, which has tracked launch performance since 2007, finds the first six months to be the critical window that sets the long-term trajectory.⁸ That is the window the Cockpit is built to watch. The Glue is the early-warning system that closes the gap between the signal and the response.

Where do you stand on the launch architecture?

The Spine · what HQ sets, every market keeps

- ☐ Do country teams use the same readiness gates regardless of the model they have picked?
- ☐ Are KPIs and dashboards identical across markets, with the same definitions and the same source data?
- ☐ Is the decision-rights matrix between HQ, the affiliate, and partners explicit, current, and known by everyone involved?

The Choice · how country teams pick the model

- ☐ Can you defend your current launch model with documented scoring evidence rather than opinion?
- ☐ Have non-viable models been eliminated with hard constraints before scoring the rest?
- ☐ Is affiliate execution capability a real constraint in the model decision, or a footnote that gets added afterwards?

The Cycle · how both move as the market does

- ☐ Are transition triggers numerical, agreed before launch, and owned by The Council?
- ☐ Does the Cockpit fire alerts on leading indicators, or only on lagging sales numbers?
- ☐ Do content libraries assemble for the next stage of the engagement journey, or only for the current one?
- ☐ When the model changes, do reps and MSLs change role at the same time, or three months later?

Six disciplines for making adoption stick

How an organisation adopts the Architecture matters as much as the Architecture itself. Six disciplines separate organisations that genuinely adopt the framework from organisations that produce a presentation about it and move on with their existing habits. All six are non-negotiable in our experience.

-
- 1 Ground model selection in validated HCP data.** Segmentation, pathway mapping, engagement readiness. Loud voices in a workshop are not data, and a model picked on opinions instead of evidence is a model that will be challenged and re-debated every quarter by people who were not in the original room.
 - 2 Build the data infrastructure before the model is chosen.** The Architecture only works when the country team running the scoring workshop has access to current HCP and market data that can validate their assumptions. Specifically, if a segmentation assumption cannot be tested against real engagement data within four weeks of launch, the model picked on top of that assumption is theoretical. The data infrastructure that produces those validations has to be running before the scoring workshop sits down, not promised for later.
 - 3 Treat Omnichannel Hybrid as the structural default.** As The Choice section argues, Model B is the right structural answer in most established categories, not a staging post on the way to Model A. The discipline at adoption is to resource it as the deliberate, long-run choice it is, and to hold that line when a more field-heavy posture starts to look more impressive than it is.
 - 4 Use market differences as a segmentation asset.** Every market is structurally different from the next one. HCP access varies, the regulatory and HTA picture varies, healthcare-system structures vary, and where partners are involved their footprints vary as well. These differences are features of the launch reality, not bugs to engineer around. The Architecture is designed to absorb them: every market picks the model that fits its conditions, inside a Spine that holds the franchise together.
 - 5 Measure HCP behaviour change, not field activity volume.** When the organisation tracks the number of calls made instead of how far HCPs have moved through the engagement journey (from awareness to consideration to prescribing confidence), the model begins to decay. Field teams optimise for whatever number they are measured on. The number that drives the model is journey progression. Activity volume is a downstream by-product, not the goal.
 - 6 Co-design the Architecture; do not roll it out.** An architecture designed in HQ and pushed down to the country teams without their input is one that the country teams will execute as a formality. When the launch pressure starts and shortcuts become tempting, a framework with no local fingerprints is the first thing that gets quietly ignored. The country workshop is where the country teams take ownership of the model and the rationale they will need to defend in the months ahead. Skipping it produces a launch that complies with the framework on paper and works around it in practice.
-

Where to start: recovering the architecture you already have

For most organisations, the Architecture is rarely built from a blank page. Most of it already exists somewhere in the building, waiting to be recovered, systematised, sharpened, and connected across the layers. The starting move is a foundation diagnostic across the three layers.

Spine

Where does the organisation already have it? Where is it fragmented across markets, with each affiliate carrying a slightly different version of what was supposed to be common? Where are affiliates actively contradicting the global strategy without anyone at HQ seeing it? Most organisations discover they already have most of The Spine, but in pieces, owned by different functions, with nobody responsible for the whole.

Choice

Which of the eight scoring dimensions does the organisation actually score with evidence? Which dimensions does it currently assume rather than test? How many markets are on the right archetype for defensible reasons, and how many are there because nobody pushed back on the original choice? When models change, do they change because of a structured decision or because of pressure from the loudest stakeholder?

Cycle

Are transition triggers documented anywhere with pre-agreed actions? Or do they live in the heads of a few senior people who may not be there for the next launch? Does the content approval process anticipate the next stage of the engagement journey, or only approve materials for the current one?

THE ARCHITECTURE COMPOUNDS WITH EVERY LAUNCH

By the second or third launch on the same architecture, the marginal cost of starting up keeps falling. The absolute investment stays high — but each launch inherits the framework, platform and country workshops the last one built, rather than work the organisation has to figure out for the first time.



APPENDIX

The Layer 02 diagnostic

The foundation is the work. Everything else is decoration.

The body of this paper introduces The Choice layer in shorthand. This appendix sets out the full diagnostic country teams use to pick the model: the four archetypes, the four hard constraints in Stage One, the eight scoring dimensions in Stage Two, and the scoring matrix — turning the model decision from a workshop opinion into a documented, defensible choice.

Overview of the four archetypes

Before the scoring workshop opens, the country team needs a shared understanding of what each archetype actually means in practice. The ‘most common misuse’ row matters in particular, because that is where most teams make expensive mistakes early in the conversation.

	A  KOL-Led Specialist	B  Omnichannel Hybrid	C  Foundation Build	D  Digital-First
CORE LOGIC	Depth over reach. Intensive, high-value relationships with a small number of the right KOLs and accounts.	Reach and frequency at scale. Field initiates the engagement and handles complexity; digital sustains, deepens, and extends it.	Build before promote. Diagnostic infrastructure, HCP education, and patient pathways are the launch; promotional activity follows once the foundation is ready.	Channel constraint turned into a strategy. Virtual, portal, and content-led engagement replace or supplement field where field is structurally constrained.
WHEN IT FITS	Concentrated centres. High therapy complexity. Reimbursement secured. Defending against fast competitive entry.	Dispersed prescriber base. Established class. Mixed channel readiness across the HCP population. The structural model for most categories.	Pre-competitive market. Unestablished diagnostic and referral pathways. Rare or novel indications. Long or fragmented diagnosis journey.	Partner-controlled channels. Niche or low-volume indication. Access restrictions in place. Digital-native HCP base.
MOST COMMON MISUSE	Applied to dispersed markets because it feels ‘premium’. The unit economics collapse: the top 200 accounts rarely drive enough volume to justify a Model A field investment.	Treated as a default or as a transitional compromise on the way to something more committed. In most established categories, Model B is the right structural model, not a fallback.	Exited too early when early KPIs are soft. Foundation Build needs the patience to do its job. Abandoning it before the pathway and education work has converted is a common, costly mistake.	Deployed as a cost-saving measure when Model A or Model B is what the market actually needs. The cost saving is real. So is the revenue miss that follows from picking the wrong model.
THE SIGNAL YOU’RE IN IT	Your top 20 accounts drive more than 60% of addressable volume. Prescriber concentration is structural, not temporary.	No single account or region dominates volume. Channel mix is the primary commercial lever rather than relationship depth.	HCPs are willing to prescribe but structurally cannot. Diagnosis or referral is the real bottleneck. Intent is not the problem.	Reaching the right prescribers requires a digital workaround regardless of how much field effort or investment is applied.

STAGE 1

Eliminate non-viable models before scoring anything else

Four constraint questions must be answered before the scoring matrix opens. The questions are binary. A 'No' answer excludes that archetype from consideration entirely. Models eliminated at Stage 1 do not proceed to Stage 2 scoring.

#	CONSTRAINT QUESTION	IF 'NO' →	WHY IT MATTERS
C1	Is reimbursement secured or imminent in this market?	Eliminates Models A and B. Default to Model C until reimbursement is resolved.	A high-touch field effort cannot be sustained when HCPs cannot prescribe. The organisation ends up building relationships it cannot yet convert, which is expensive goodwill with no commercial return.
C2	Does the organisation control its own field deployment in this market?	Eliminates Model A. Constrains Model B. May force Model D or a partner-shaped model.	A co-promote or distribution partner with its own field priorities will erode any model that depends on flexible field deployment. Alignment on paper does not survive competing brand pressures in the field.
C3	Is diagnostic infrastructure established in this market (pathways, testing, referral routes)?	Eliminates Models A and B. Forces Model C first.	Field effort deployed into a market where patients cannot be identified or referred is expensive education with no short-term conversion pathway. The launch has to build the diagnostic capacity before it can build prescribing volume.
C4	Can the affiliate realistically execute the model that the scoring will point to?	Downscale to what the affiliate can actually deliver. Build capability in parallel for the future model.	This is the most consistently ignored constraint in global launch planning. The right model for the market, run badly by an underpowered affiliate, produces worse outcomes than a simpler model executed well.

STAGE 2

Score eight dimensions; the dominant column reveals the model

Once Stage 1 has eliminated the non-viable models, eight scoring dimensions shape The Choice among those that remain. The dimensions are not equal in weight. Some are stronger signals than others in specific therapy areas and market contexts. Each one surfaces something the other seven cannot.

#	DIMENSION	WHAT IT MEASURES	WHY IT DRIVES MODEL CHOICE
1	Prescriber base structure	How concentrated is volume across accounts and HCPs in this market?	Concentrated bases justify deep, resource-intensive field investment. Dispersed bases need scalable channel mix instead. Model A's economics collapse when the target list has thousands of names.
2	Competitive intensity and strategic posture	Are you defending an established position, attacking a competitor, or building before competitors arrive?	Defensive postures need Model A's depth. Offensive postures in high-volume markets need Model B's reach. Pre-competitive windows demand Model C's infrastructure focus before promotional noise becomes useful.
3	Value-prop complexity and narrative maturity	How much clinical and conceptual work is needed to shift prescriber behaviour?	A novel mechanism requires shaped conversations only Model A can sustain. An established-class entrant needs reach and consistency, which is Model B's territory. An undereducated market needs Model C's foundation work before the value proposition lands.
4	Market and HCP accessibility	Can you reach the right prescribers, and through which channels will they actually engage?	Folds healthcare-system readiness and HCP channel preference into one signal. High physical access and high digital readiness opens Models A or B. Constrained physical access with viable digital reach opens Model D. Neither at scale points to Model C.
5	Access, reimbursement, and HTA landscape	What is the reimbursement status, and how does payer or HTA pressure influence prescriber autonomy?	Where a formal HTA or payer-pricing process gates access, prescriber willingness does not translate into volume. High payer control compresses Model A's upside and reshapes Model B's channel economics. Pre-reimbursement conditions force Model C.
6	Channel ownership and partner dynamics	Who controls the channels to market, and do partners' priorities align with your launch?	The most underrated dimension in launch planning. A co-promote partner with competing brand priorities changes every assumption a model rests on. In extreme cases this becomes a hard constraint (see C2 above).
7	Patient journey and diagnostic infrastructure	How established are the pathways from symptom to diagnosis to treatment initiation?	A fragmented diagnostic pathway means even willing KOLs cannot act on their intent. Field effort spent in front of HCPs whose patients cannot be identified produces goodwill rather than prescriptions.
8	Medical affairs lead-time versus commercial window	How far ahead of commercial launch has medical built the scientific foundation in this market?	Model A and Model B can move at pace when medical is well ahead. When medical and commercial launch in parallel, Model C's slower cadence is the more honest acknowledgement of what is achievable.

The full scoring matrix

For each scoring dimension, read all four archetype descriptions and identify which one best describes the market the team is launching into. The dominant archetype column across all eight dimensions points to the model that fits. Where signals genuinely conflict, the conversation that surfaces in the workshop is more valuable than the tally. Amber note rows carry qualifications to read before scoring the dimension they sit under.

A KOL-Led Specialist	B Omnichannel Hybrid	C Foundation Build	D Digital-First
Dimension 1 • Prescriber base structure			
Highly concentrated. The top 10–30 accounts or KOLs drive the majority of volume. Relationship depth with a small universe is the lever.	Dispersed across a broad base. No single account dominates. Scale and channel mix matter more than relationship depth.	Fragmented or underdeveloped. The prescriber base is still forming: specialists are few, pathways are immature, or the indication has not yet embedded in routine clinical practice.	Niche or access-restricted. The prescriber universe is small or geographically dispersed in ways that make physical field coverage economically unviable.
Dimension 2 • Competitive intensity and strategic posture			
Defensive. Late entry into established territory, or protecting a threatened position. KOL depth is the moat.	Offensive. High-volume market where reach and frequency are the advantage. Or an established class where broad engagement sustains share over the long run.	Pre-competitive or first-mover. Competition is not yet the primary challenge. Field effort builds the market rather than the brand.	Constrained share of voice. Channel restrictions, partner dynamics, or budget realities limit competitive posture regardless of intent.
Note. Where competitive intensity is very high and you are defending a concentrated account base, score A. Where you are attacking a dispersed market or competing in an established class at scale, score B.			
Dimension 3 • Value-prop complexity and narrative maturity			
High complexity, narrative still forming. Novel mechanism or category-defining claim that requires shaped, two-way conversation with expert advocates.	Moderate complexity, established class. The value proposition is understood. The task is consistency, reach, and reinforcement rather than shaping new thinking.	High complexity, but the infrastructure is not yet ready to receive it. HCPs cannot yet act on the value proposition because diagnosis, pathways, or reimbursement are unestablished.	Low to moderate complexity. The therapy is well-understood. The engagement task is access and convenience rather than clinical shaping.
Note. A complex value proposition is different from a Foundation Build requirement. Model C is about market readiness to act, not about message complexity. A novel therapy in a market with an established diagnostic pathway and secured reimbursement belongs in A or B.			
Dimension 4 • Market and HCP accessibility			
High physical access to a concentrated target group. KOLs and key accounts are accessible and engagement-willing. Digital is supplementary rather than structural.	Mixed channel readiness across a broad base. Some HCPs prefer digital. Others require field contact for initiation. The model runs both simultaneously by design.	Low access because the infrastructure that would let HCPs act on the launch does not yet exist in the market. This is a system-readiness issue, not an HCP-preference issue.	Low physical HCP access, but a viable digital channel still reaches the prescriber base. Geographic, institutional, or access-policy barriers make field deployment impractical, while the HCP population can be engaged through digital means.
Note. Folds healthcare-system readiness and HCP channel preferences into one signal. High physical access and high digital readiness opens Models A or B. Constrained physical access with viable digital reach opens Model D. Neither physical nor digital reach at scale points to Model C, where the work is to build the access infrastructure itself.			

A KOL-Led Specialist	B Omnichannel Hybrid	C Foundation Build	D Digital-First
Dimension 5 · Access, reimbursement, and HTA landscape			
Reimbursement is secured or close to secured. The payer environment does not heavily restrict prescriber autonomy. KOL advocacy translates directly into prescribing behaviour.	Reimbursement is secured. Payer influence on prescriber behaviour is present but manageable. Omnichannel reach supports both clinical and access-related messaging.	Pre-reimbursement, or reimbursement contested or delayed. Field investment now is infrastructure-building rather than volume-generating.	Variable or access-constrained. Reimbursement may exist, but access barriers (formulary restriction, step therapy, prior authorisation) suppress volume regardless of HCP intent.
Note. Payer and HTA influence on prescriber autonomy sits explicitly within this dimension. In markets where a health-technology-assessment body or national payer controls pricing and funding, prescriber willingness to advocate for a therapy does not translate into volume the way it does in less payer-controlled markets. The effect compresses Model A's upside and reshapes Model B's channel economics.			
Dimension 6 · Channel ownership and partner dynamics			
Fully internally controlled. Field deployment, channel mix, and targeting are organisational decisions without partner dependency.	Hybrid control. The organisation owns the core field; partners contribute in specific channels, geographies, or segments. Partner alignment is manageable.	A field partner is available and aligned for the foundation work. Partner incentives align with pathway-building and HCP education, which suits a pre-commercial market.	Channel-restricted. The organisation does not control primary access to HCPs. Partner field allocation, distribution agreements, or access policies determine what is actually reachable.
Note. The most underrated dimension in launch planning. A co-promote partner with competing brand priorities will erode any model that depends on deployment flexibility. Stress-test the partner's incentive structure before scoring: are they rewarded for outcomes that align with your model, or for outcomes that conflict with it?			
Dimension 7 · Patient journey and diagnostic infrastructure			
The diagnosis pathway is established and functioning. HCPs in key accounts can identify the right patients and initiate treatment without structural friction.	The pathway is established across a broad base. The diagnosis rate and referral infrastructure are consistent enough that a dispersed prescriber base can act on HCP intent at scale.	The pathway is underdeveloped, fragmented, or absent. Patients exist but are not reaching the right specialists, or are not being diagnosed at all. Field effort has to build the diagnostic and referral capacity first.	Niche or low-volume indication. The pathway may be established, but patient numbers are small enough that building dedicated field infrastructure is economically disproportionate to the opportunity.
Note. Arguably the most consequential dimension for rare-disease and specialist-category launches. A field team in front of willing KOLs in a market where patients cannot be found or routed produces goodwill rather than prescriptions. If the answer here is 'pathway absent', Constraint C3 in Stage 1 should already have directed you to Model C.			
Dimension 8 · Medical affairs lead-time versus commercial window			
Medical is well ahead of commercial. Publications, guideline involvement, KOL relationships, and the scientific platform are established before commercial launch. Model A can move at pace because the scientific foundation is in place.	Medical is ahead or running in parallel, while the commercial model is designed for breadth rather than depth. Omnichannel reach does not depend on deep pre-established KOL relationships in the same way Model A does.	Medical and commercial are launching in parallel, or medical is behind. Foundation Build's slower cadence is the honest acknowledgement that commercial promotion is outrunning the scientific and infrastructure readiness in this market.	Medical maturity is less critical because field deployment is constrained regardless. Digital-First compensates for access limitations rather than depending on medical groundwork.
Note. An internal-readiness question rather than a market-characteristics question, which is why it sits last in the scoring. A launch team that picks Model A while its medical affairs function is six months behind has not chosen a launch model so much as set itself a target it cannot yet hit.			

What the scoring does, and what it cannot do

Three honest limitations of any scoring framework, and this one in particular.

1 Judgment still matters

The tool structures the conversation and makes the rationale explicit on paper. It does not substitute for experienced commercial and medical judgment about which signals matter most in a specific therapy area or market context.

2 The score reflects a moment, not a future

Scores describe the market as it is at the point of scoring. This is precisely why Layer 03, the Cycle, exists: to design the transitions before they are needed, rather than after the market has moved.

3 Resourcing is a separate conversation

Model selection and resource allocation are related but distinct decisions. Portfolio overlap and customer-overlap considerations, deliberately not in this scoring matrix, are critical inputs to the resourcing conversation that follows from the model decision.

Sources, acknowledgments, and authors

Sources

- 1 Deloitte Center for Health Solutions. 'Key factors to improve drug launches' (2019). Of launches that missed their forecast, 50% had limited market access, 46% reflected an inadequate understanding of local physician and patient needs, and 44% had weak product differentiation.
- 2 Trinity Life Sciences. 'Moving the Needle: Lessons from the 2023 Launch Class' (November 2024). 50% of the 2023 US launch class underperformed its first-year forecast, an improvement on 54% across 2020–2023.
- 3 Bristol Myers Squibb, fourth-quarter and full-year 2025 results. Sotyktu worldwide revenues of \$291 million in 2025, up 19% on 2024.
- 4 NHS England, letter on cardiovascular disease prevention (31 March 2025), reported in The Pharmaceutical Journal, 'NHS increases funding for inclisiran although uptake remains "slower than expected"' (3 April 2025). GlobalData, Pharma Intelligence Center analyst consensus forecast (February 2022): Leqvio projected to reach \$2.5 billion in sales by 2027. Novartis full-year 2024 results: Leqvio global net sales of \$223 million.
- 5 Lublóy Á. 'Factors affecting the uptake of new medicines: a systematic literature review.' BMC Health Services Research, 2014. Prescribing habit is among the strongest predictors of new-drug uptake; adoption typically plateaus within roughly six months of launch.
- 6 Deloitte Center for Health Solutions. 'Drug launches reflect overall company performance' (analysis of 284 US launches, 2012–2021). Of drugs that miss expectations at launch, only ~26% reverse trajectory in year two or three, and 18% reverse it durably across both.
- 7 McKinsey & Company. 'The secret of successful drug launches' (2014). Approximately two-thirds of drug launches miss pre-launch expectations.
- 8 IQVIA, Launch Excellence series (published since 2007). The first six months of launch is the critical period that determines long-term success.
- 9 Fractal Force analysis of the 50 most-anticipated drug launches of 2020–2025, setting pre-launch analyst consensus or company guidance against company-reported actuals. Of the 41 with a clear pre-launch peak forecast, 17 delivered under 40% of it; of the 15 clear misses, 2 were clinical or regulatory failures. COVID products excluded; figures on mixed bases and FX, some 2025 values estimated.

Acknowledgments

This paper draws on work and conversations with commercial, medical, and cross-functional leaders across global and regional life-sciences organisations between 2024 and 2026, on Fractal Force's own analysis of recent launch performance, on published research from McKinsey, IQVIA, ZS Associates, and Deloitte, and on peer-reviewed work on physician adoption of new medicines. Any errors of interpretation are ours.

— ABOUT FRACTAL FORCE

The foundations that decide whether HCP engagement actually works.

Fractal Force is a life sciences consultancy focused on one thing: the commercial and medical foundations that decide whether HCP engagement actually works. We help companies move from a brand-centric model to a customer-first one across five connected areas — data and analytics, segmentation, engagement journeys, operating models and impact measurement — and make sure that omnichannel, CX, AI and other investments already in place deliver ROI.

We work with the commercial, medical, digital, and brand leaders at mid-to-large pharma, biotech, and medtech companies.

We deploy a global, plug-and-play network of experts and tech partners that spans strategy through implementation. Unlike traditional consultancies and agencies, our model keeps us lean, fast, and competitively priced. Senior people only — no bench to keep busy, no juniors on your budget, no house method bent to your problem. We architect the strategy and stick around to build it. That's the last mile the big firms skip, and the implementation shops can't reach.

WORKING ON A LAUNCH?

We help global and country teams build the **Launch Decision Architecture**.



Architects of customer-first foundations for life sciences companies.

CONTACT

✉ info@fractal-force.com

in LinkedIn

☎ +32 470 865302

