

Endogenius™ Precision Cancer



1) What Endogenius™ is

Endogenius™ is a **patent-pending AI framework** that provides medical teams with a **decision-support engine** — not a consumer chatbot and **not an LLM** (ChatGPT, Gemini, Claude). LLMs invent fluent answers that are hard to audit in a tumour board. We do not play that game.

Endogenius™ is based on theoretical work in intelligence, as outlined in the 2018 book *What's on Their Mind?* It integrates a computable engineered model based on **Friston's free-energy principle** (Bayesian brain) and **Hawkins' memory-prediction**. These meet in perception module **A**, where memory forecasts the next step of a sequence, and free energy measures surprise when reality disagrees. As intelligence needs more than perception, the framework integrates an action module **B** and a control loop module **M** (feedback under constraint).

2) $F(T) = (A \rightarrow B)$ under M

This simple formula that we adapt across industries and apply in healthcare / pharmacology (precision oncology, rare-disease diagnostic, geriatric pharmacology, drug combo) integrates **three modules and one combined formula**:

Module / Combined formula	Role	Example
A Perception	Structure short- and long-term memories into a knowledge library; fuse unsupervised + supervised algorithms (weighted) to identify patient-specific cancer patterns.	Breast cancer + chest RT (radiotherapy / radiation treatment to the chest) + host risk → later-cancer pattern; SEER-like pathways flag lung / other breast.
B Action	Takes Module A's pattern findings and turns them into concrete options for the board: what to call the case, which drug cocktail / dose ideas, and which visit window.	"Treat as likely new lung primary — propose lung-adapted drugs and checks in years 1–5; do not copy the old breast protocol."
M Control	Gate that decides whether the action is allowed. Oncology golden rules originating in: IKNL (Netherlands), NCCN, ESMO, SEER MPH / IARC multiple-primary, ECIS mortality priors.	Missing histology, metastasis vs second-primary conflict, or unsafe combo → hold .
F(T) + BDN (combined formula)	F(T) = ranking score for the options the board can compare (audit trail). BDN (Bayesian Dynamic Network) = separate trust score: how confident we are that the evidence is solid enough to show that ranking.	Breast survivor, new lung lesion: F(T) may put "new lung primary" at 0.68 and "breast metastasis to lung" lower. BDN answers a different question: is the work-up complete enough to trust that ranking? No histology yet → BDN low → Module M holds with a written reason. The 0.68 is for the board's file, not an automatic order.

In practice we run **three F(T) scores** on the same spine — origin (CUP), organ, and body — then Module M decides whether any action may pass:

F(T) type	What it answers	Example
F(T) origin (CUP)	Where did this cancer start? Primary versus metastasis when the origin is unclear (cancer of unknown primary).	A nodule with no clear primary: rank likely tissues of origin from gene-activity patterns — or refuse when trust is too low.
F(T) organ	How each involved organ changes the plan — dose, toxicity, and local constraints.	Lung lesion after breast cancer: organ score flags lung-specific toxicity and visit windows, not a copy of the old breast protocol.
F(T) body	Whole-patient host safety: comorbidities, prior therapy load, and later-cancer foresight across the body.	Same woman: body score asks whether the proposed cocktail is survivable given kidney function, prior chest RT, and her later-cancer watch-list.

The practical objective is a **cocktail of drugs with the right dosage earmarked for each organ**, while taking into account the patient's **exogenous factors** (prior radiotherapy and chemotherapy, kidney and liver function, other diseases, age, and host risk) — then only release that proposal when BDN trust and Module M allow it.

3) The patient's journey

Each patient has a history: index cancer, therapy, host, molecular clues. We integrate that history by interfacing to existing hospital / clinician data sources. From that history, we expect how the case should look next. When it stops fitting — metastasis versus a new primary, a CUP that does not match, foresight evidence too thin — we do not push a fluent answer. **We hold**.

Clinical questions the engine clarifies

Endogenius™ is not limited to follow-up risk. The same spine covers label (primary vs metastasis vs unknown origin), organ-adapted drug cocktail ideas, and survivor foresight — or **holds** when trust is too low. Running example: a woman treated for breast cancer years ago now has a new lung lesion.

Situation	What staff ask	What the tool aims to show
After a first cancer	Same disease, or a new one?	Metastasis vs independent second primary (secondary cancer)
Nodule / CUP	Where did it start — primary, or metastasis from elsewhere?	Tissue-of-origin hypotheses (F(T) origin); refuse when unsure
Known metastasis / multi-organ disease	Which drugs, at what dose idea, for which organ — safely?	Systemic + organ-adapted cocktail proposals (F(T) organ + body), under golden-rule / BDN gates
Survivor under follow-up	After radio / chemo — what to watch, where, and when? (not a generic "survivors need CT" — CT = computed tomography scan)	Later-cancer risk band, organ watch-list, 1 / 5 / 10-year windows for tumour-board review

Before any of those outputs reach the board, Module M still runs four **gates** on the case: (1) what should we expect for *this* patient; (2) is evidence solid enough today; (3) is the suggested plan safe for her; (4) are we allowed to recommend — or **hold** with an auditable reason.

Closing contrast. Many AI tools are built to always give an answer — a label, a probability, a plan — so the conversation keeps going. Endogenius™ is built to **stop** when the case does not fit the evidence. We would rather say “**we cannot recommend yet**” than produce a polished 72% that looks scientific but rests on thin data. The hard work is deciding when to stay silent and which plans must not be offered — not making the score look better on a chart.

4) What our numbers already show (Sept 2026)

Survivorship foresight

Question	In plain English	What we measured (RUO)
Whether	Will this survivor get another cancer? We show a risk band (higher / mid / lower) — not “you have a 37% chance.”	On ~100,000 SEER patients (later years held out), ranking by risk scored AUC ≈ 0.74 . Useful for banding groups; not a sharp personal probability.
Where	Which organs should the board watch first? We show a short organ watch-list, not a certainty.	True later site fell in our top-3 about 47% of the time (hand-made organ lists were ~28%).
When	Over what horizon? We use 1-, 5-, and 10-year windows — not “year 7 at 40%.”	Timing discrimination (Fine–Gray C) ≈ 0.58–0.61 — time bands, not calendar precision.

Unknown primary (CUP)

What we tested	Result	Meaning
Tumours where doctors already knew the organ; gene-activity numbers only	≈ 97–98%	Not hospital CUP yet
Hard practice test (noise + hidden genes)	≈ 61% (bar 60%)	Refuse when unsure
Real hospital CUP RNA (Melbourne + Heidelberg, locked)	Pending	Needed before “we found the primary”

5) What we are seeking

Priority. Name a **pilot centre** and a small first set of real cases, so we can sit together through one example report (drug combo and visit plan) — research use only. **That leads to** anonymized follow-up (first cancer, treatment, later cancer site and date) to recalibrate later-cancer numbers for your population. **To succeed we need to agree** how we speak about results: a risk band, not a personal percentage; organs to watch, not a certain diagnosis; when to act versus when to hold.

Contact: serge.vanthemsche@endogenius.ai · Endogenius.ai / PM&BD; Consulting · Health