

AI GOVERNANCE

The Use Case Statement

Governing what a tool **does** —
not what it is.

A one-sentence template for AI committees. Copy it,
change it, make it yours.

THE PROBLEM

“We've approved 40 AI tools” is a list, not an inventory.

Ask a health system what its AI exposure looks like and you'll usually get a spreadsheet of product names.

Which means when someone asks “**what's our risk on ambient documentation?**” — the answer is one row.

But the tool isn't the thing you approved. The **deployment** is.

WHY IT MATTERS

One tool. Three deployments. Three risk objects.

SAME PRODUCT

Ambient documentation, primary care. Clinician reviews and signs every note before it enters the record.

Ambient documentation, emergency department. Autopopulation on, review after the fact.

Ambient documentation, behavioral health. Recording in a session covered by additional state consent rules.

DIFFERENT RISK

Human in the loop. Advisory impact. Low exposure.

Human on the loop at best. Determinative content in the legal record. Materially higher.

A different statute applies entirely. Some states now require explicit consent for recording.

The committee that approved the first did not approve the other two.

THE TEMPLATE

Write every use case as one sentence.

THE STATEMENT

[Product] [function] [object] from [input] to produce [impact] [output] affecting [population], with [human control], in order to [objective].

It reads awkwardly on purpose. The awkwardness is the point — you cannot complete the sentence without saying out loud who it affects and how much human control there is.

Most governance failures we see aren't disagreements about risk. They're the committee and the implementer having different sentences in their heads.

Some slots are a vocabulary. Some are free text. That's deliberate.

GOVERNED — PICK FROM A LIST

Function (multi) — comparability. You need to find every tool in the portfolio that **decides**.

Input (multi) — data lineage. Drives most privacy and boundary questions.

Impact (single) — the most consequential field on the form.

Population (multi) — who bears the risk.

Human control (single) — the mitigating variable in almost every risk you'll write.

FREE TEXT — SAY IT IN YOUR WORDS

Object — what it acts on. "Encounter audio," not a category.

Output — what actually lands in the workflow.

Objective — the business justification, in the sponsor's own words.

Product is static: the vendor's tool, one record, many use cases hanging off it.

Governed slots make the portfolio queryable. Free-text slots make two deployments distinguishable. You need both.

THE VOCABULARIES

Four lists do most of the work

FUNCTION

classifies · scores or ranks ·
extracts · predicts · generates ·
recommends · summarizes ·
detects · **decides** · automates

IMPACT

advisory — informs a human
who decides

influential — shapes what the
human sees or considers

determinative — the output is
the decision

INPUT

EHR · clinical notes · imaging ·
vitals · medications · lab results
· problem list · claims · coding
data · eligibility & payer data ·
voice

HUMAN CONTROL

in the loop — human acts on
every output

on the loop — human
supervises, can intervene

advisory only — no action
taken from output

autonomous — no human in
the path

The same product, written twice

DEPLOYMENT A

[Ambient scribe] **summarizes** the clinical encounter from **voice and EHR context** to produce an **advisory** draft progress note affecting **patients and providers**, with a human **in the loop**, in order to reduce documentation burden in primary care.

DEPLOYMENT B

[Ambient scribe] **generates** the clinical encounter record from **voice and EHR context** to produce a **determinative** signed note affecting **patients and providers**, with a human **on the loop**, in order to reduce time-to-chart-closure in the emergency department.

Two sentences. Two risk profiles. Two sets of controls.
One product name.

Why not make everything a dropdown?

Because if every slot is governed, you get a clean database and a useless one. Every ambient tool collapses into the same row and you've lost the thing you needed to govern.

The vocabularies exist so you can **query** the portfolio — show me everything determinative that touches patients with no human in the loop. The free text exists so a reviewer can tell two deployments apart.

ONE RULE THAT MATTERS

Never reject someone's added vocabulary term. Let them add it, then clean it up later. A governance system that argues with people at intake doesn't get used.

Three steps, in order

- 1 Rewrite your five highest-risk existing approvals as statements.

Do the ones you already approved. You will find at least one where the sentence you write today doesn't match what the committee thought it approved. That finding is the business case.

- 2 Make the statement a required field at intake.

Drafted by the tool sponsor, not the committee. If the sponsor can't complete the sentence, the submission isn't ready — and you've saved a meeting.

- 3 Make it the unit everything else hangs off.

Risks, controls, monitoring, and the final report all reference the statement. Edit it in one place and it updates everywhere, or you'll have four versions inside a month.

Expect step 1 to be uncomfortable. That's the signal it's working.

Steal this. I'd rather it be useful than cited.

This is the template we use, developed across roughly thirty health systems. The vocabularies below are a starting point, not a standard — every system's are slightly different, and they should be.

Troy Bannister — Founder & CEO, Onboard AI

If you adapt it and something breaks, I'd like to hear about it. That's how it gets better.