



Embedding AI into Clinical Trial Workflows

UX Design for Intelligent, Compliant
Systems



Narendra Parmar
Sycamore Informatics

phuse.global

connect·share·advance









The Law of Conservation of Complexity

“Every application has an inherent amount of irreducible complexity. The only question is: Who will have to deal with it, the user or the product team?”



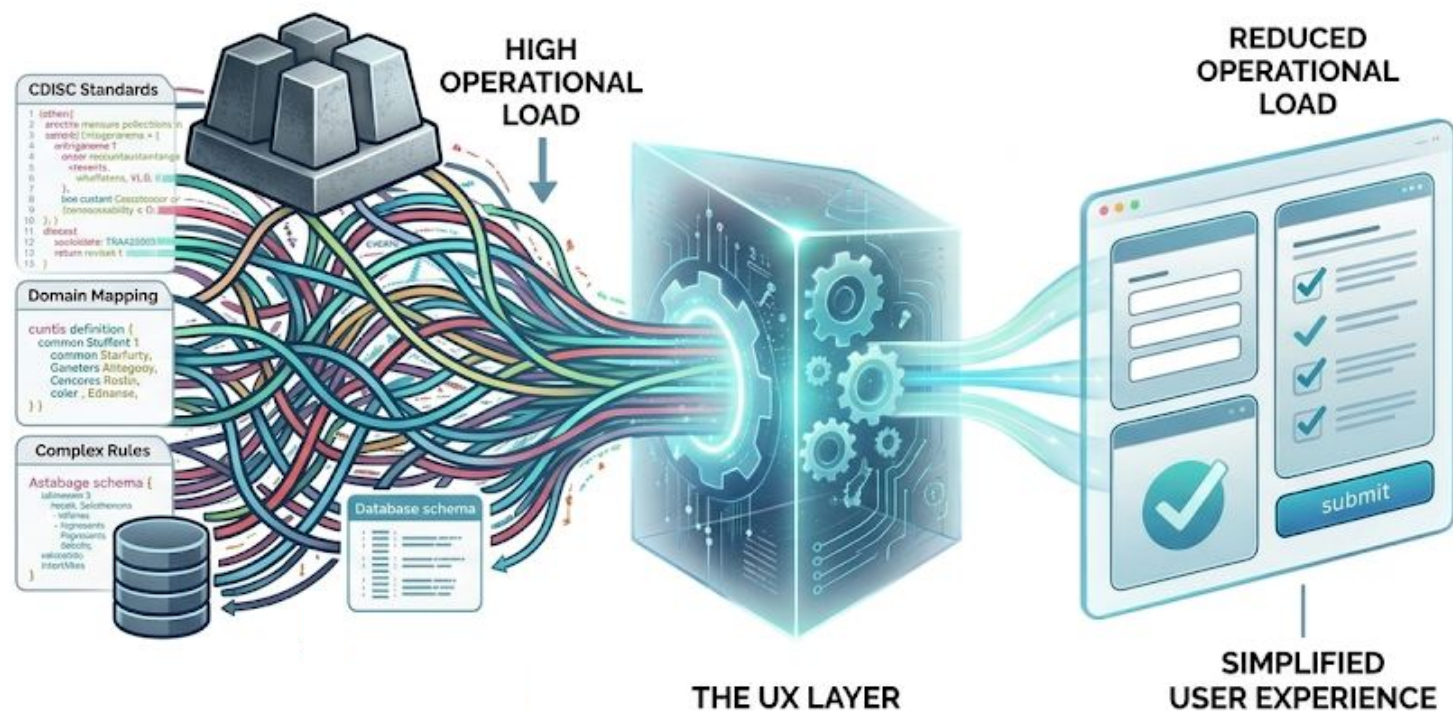
Larry Tesler

Computer Scientist & Pioneer of Interaction Design



Our Philosophy

Shift the complexity from the ~~user~~...
...to the product.



WITHOUT UX

Users must learn the software.

Cognitive load falls on the programmer.



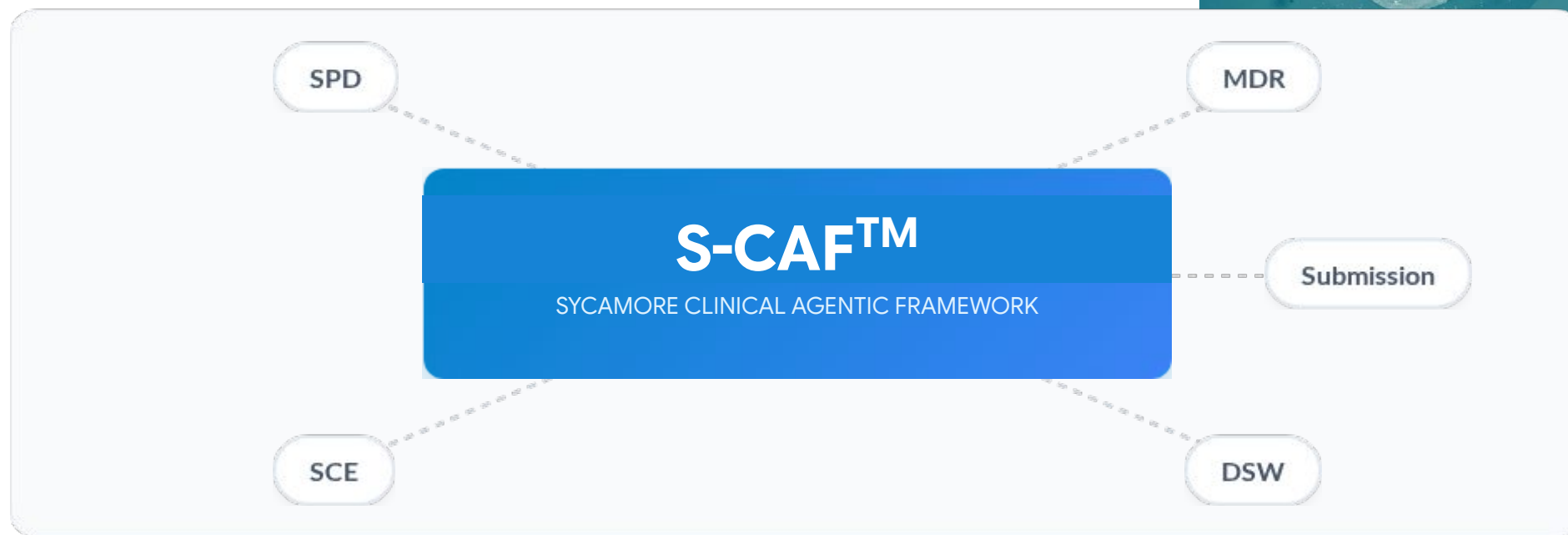
WITH UX

Software understands the user.

Operational load is absorbed by the tool.

phuse Designing for the Agentic Future

Google Workspace has Gemini. Microsoft has Copilot. At Sycamore, we are building **S-CAF™**.



AI only succeeds when users trust it. Good UX makes AI explainable.

Users should always understand WHY, WHAT, and WHERE before accepting an AI recommendation.



connect·share·advance

Why, What & Where

Key Patterns to use AI in Clinical Data Workflows



The Trust Deficit

01 . COMPLIANCE & INTELLIGENCE

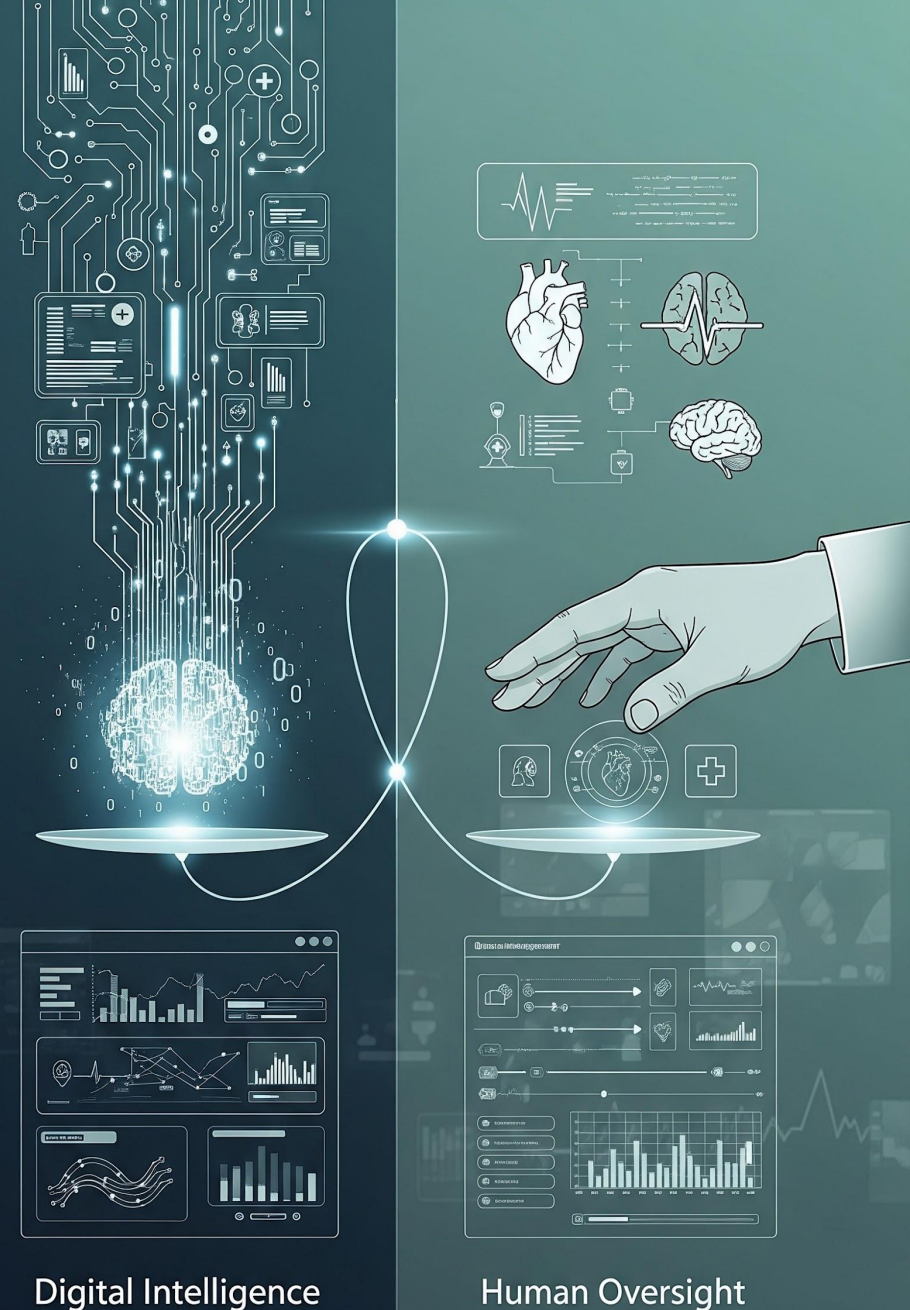
Clinical software demands both sophisticated intelligence and strict regulatory compliance.

02 . USER CONTROL

Data managers and regulatory specialists need AI features that feel transparent, controllable, and inherently safe—not black-box automation.

03 . THE TRUST PARADOX

The core design challenge: The users who need AI the most are trained to trust it the least.





Showing a Visual mapping





Showing a Visual mapping and Contextual information
for users to reduce decision time



VARIABLE MAPPING

PROGRAMMER VIEW

AGE → DM.AGE

SOURCE · FORM ITEM

FormDM

OIDLP.DM.AGE

Data Typeinteger

Codelist—

TARGET · SDTM VARIABLE

DomainDM

VariableAGE

LabelAge

CoreExp

MAPPING FUNCTION

XCOPY(DM.AGE)

TEST DATA PREVIEW

3 SAMPLE ROWS

SUBJID	DM.AGE (RAW)	→ DM.AGE
S-001	42	42
S-002	67	67
S-003	29	29

Preview from most recent regen

Run against latest snapshot

DOWNSTREAM IMPACT

3 CONSUMERS

ADSL

ADAE

TLF T14.2 (Demographics)

ORIGIN & RATIONALE

STANDARD

Direct standard mapping — AGE is a CDISC SDTM DM special-purpose variable.

REVIEW DECISION

Current:

APPROVED

Add rationale / feedback (optional)

Approve

Reject



Pattern 1: Transparent Entry Points

01 No Ambient Automation

If an AI mapping quietly changes, the audit trail is broken. Automation must be explicitly triggered.

02 Visible Separation

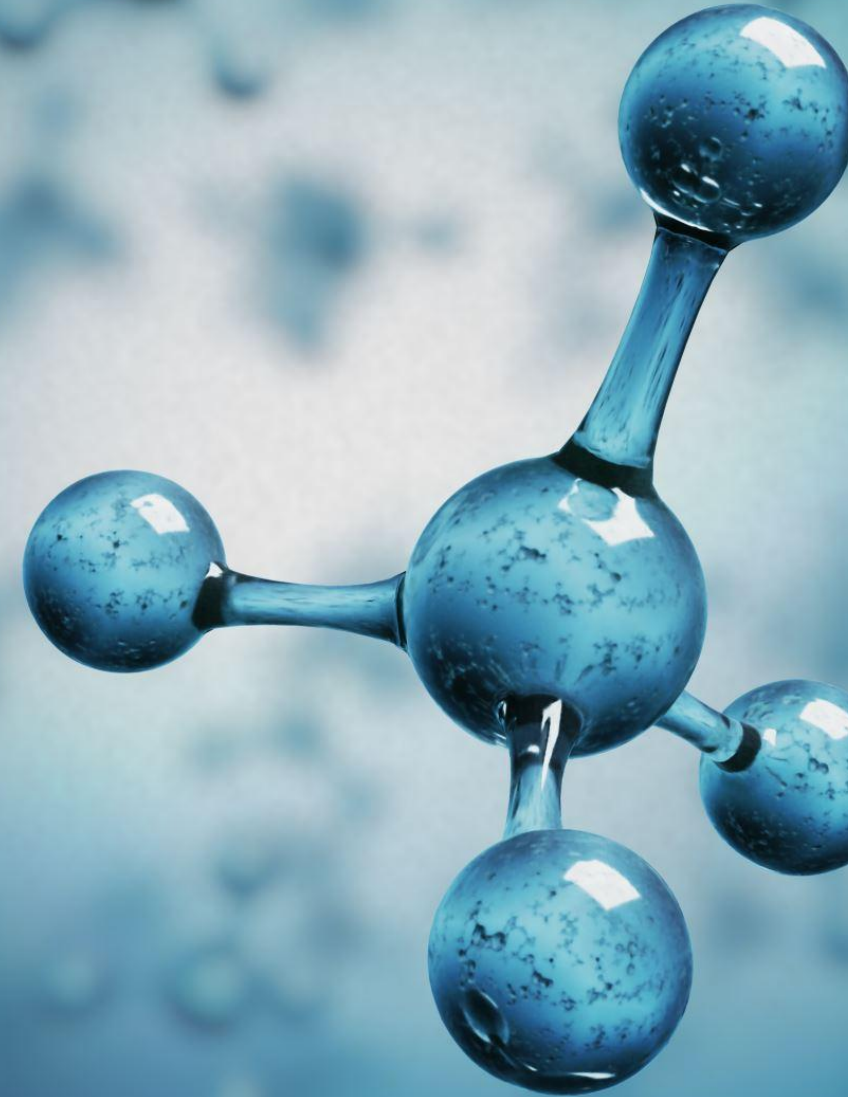
We built "Suggested Mapping" panels that sit parallel to human mapping tables, never overwriting records automatically.

03 Shared UI Languages

AI-assisted flags reuse the same visual patterns as human-authored flags.

04 Predictability

The user's existing mental model of "what happens when I see this" remains entirely intact.





No Ambient Automation



- 👁 "Recommend Compute" proposes a resource change in a **side-by-side panel** — it never writes into the exec config until the user accepts.
- 👁 Two visual states, always distinguishable: **suggestion** vs. **committed**. Never a moment where you can't tell which one you're looking at.



Same Flag, no new Visual Language

SCE · 00001rtm_project_fya — files		
NAME	STATUS	AI INSIGHTS
 SIM_CMP_0512.sas7bdat	SAS Dataset	3 unmapped vars
 define.xml	1 Error	Validation error
 validation_report.pdf	Stale	6 months old

- Same status badge, same inline chip in the row — we didn't invent an "AI found this" visual language.
- Introducing a second, AI-specific look would have implied AI findings were a **different category of trustworthy** — we didn't want that implication either way.



Pattern 2: Explainability & Confidence

Confidence Scoring

A simple visible indicator shows how confident the system is in its suggestion. This empowers data managers to triage effectively:

high-confidence suggestions get a quick glance-and-accept, while low-confidence items receive thorough human scrutiny. Attention is spent where it is truly needed.

Explainable Rationale

Every AI output carries its reasoning, not just its conclusion. Explanations like "matched on variable name + terminology overlap" allow users to verify the logic. In a regulated domain, "trust me" is not a valid UI pattern; evidence is required for adoption.

phuse Confidence you can act on

MDR · AI Review — VS domain

VS — VITAL SIGNS

→ BP_SYSTOLIC → VS.VSORRES VLM: VSTESTCD = SYSBP High ✓ ✕

→ VSTESTCODE → VSTESTCD Multi-source: also fed by VS_7 Medium ✓ ✕

LB — LABORATORY

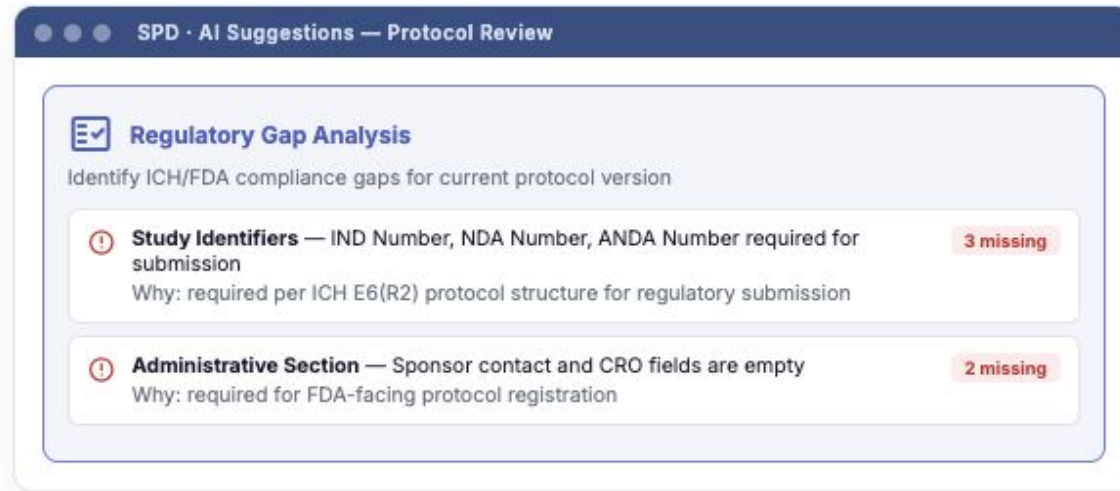
→ TESTCODE → LBTESTCD Matches CT version + 12 prior studies High ✓ ✕

High 75-100 Medium 26-74 Low 1-25

- Every suggestion now carries a **High / Medium / Low** signal, so attention goes where it's earned.
- Paired with a specific rationale — **CT version match, prior study count** — in the language a standards-trained reviewer already thinks in.



Naming the gap, not just flagging one



- ✓ Regulatory Gap Analysis aligns findings to **ICH E6(R2)** and FDA guidance, not a generic completeness score.
- ⚡ A standards lead verifies the gap in seconds instead of re-deriving it from the protocol.



Pattern 3: Proactive Detection

Proactive Flags

The system surfaces likely issues early, front-loading discovery before downstream processes fail.

Compliant Log

Human actions and accepted AI suggestions are logged identically, maintaining absolute regulatory control.

Continuous Run

AI-driven validation runs continuously against incoming data, rather than waiting for a manual batch check.

Human Decision

Every flag opens a review surface. The human decides to accept, edit, or dismiss. AI never auto-resolves.

DSW · Exec Config — test45PN · Diagnose Build

✓ Base Image Pull · rocky:linux

✓ OS Libraries · Installed

✗ R Package Install — FAILED on openssl@2.1.1
Compilation error: requires libssl-dev ≥ 3.0, base image provides 1.1.1. Package install aborted before tidyverse/dplyr/htrr could complete.

Root Cause Analysis — Suggested Action

openssl@2.1.1 needs a newer libssl-dev than the pinned rocky:linux base image ships. Quick fixes: pin openssl to 2.0.6 (matches current libssl-dev), or upgrade the base image to pick up libssl-dev 3.0+.

Apply Fix

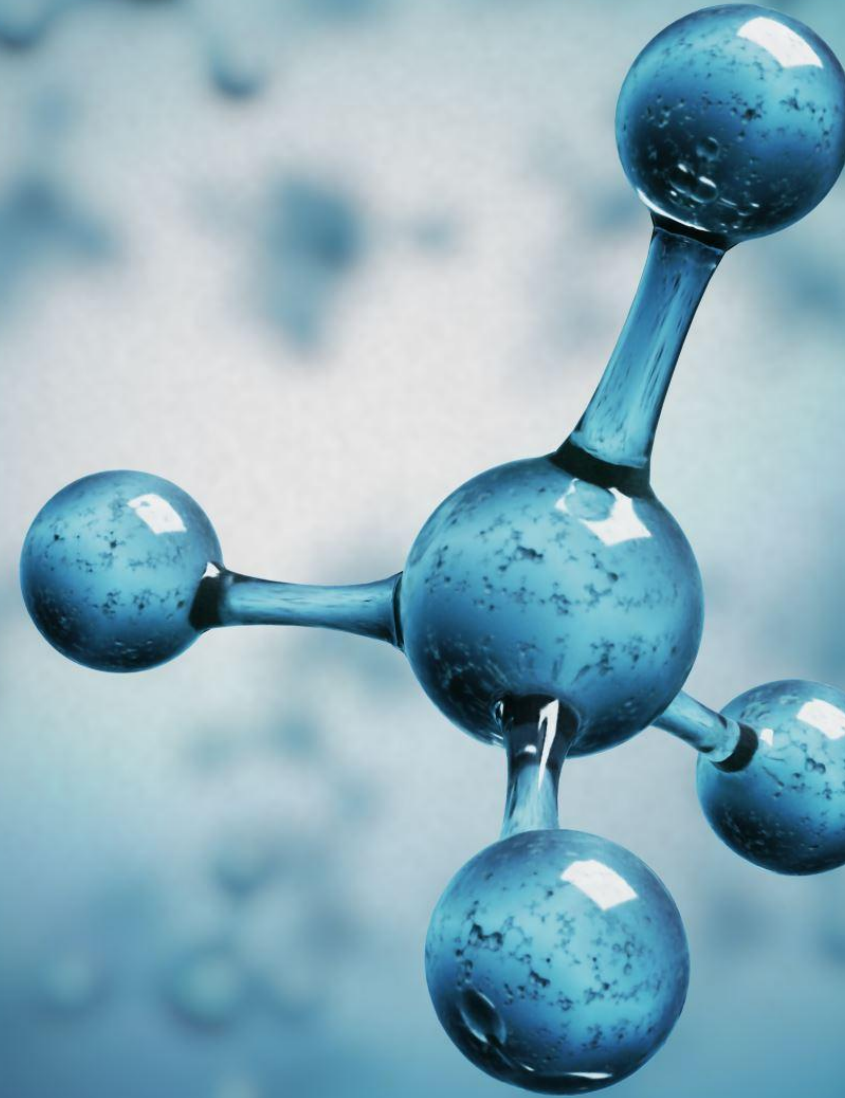
Edit Dockerfile

Proactive detection **never auto-resolves**. Every flag opens into a review surface: accept, edit, dismiss, or escalate.



The Core Design Principle

“Where does a human make a judgment call, and where can a machine make a suggestion without making a decision?”





UX laws that hold up in AI

01

System Status

Users must always know whether they're looking at AI output, human input, or a suggestion pending action. Hide the source, and trust collapses.

02

Language Match

Match system language to user language. "High confidence" means nothing to a data manager. "Matches CT version + 12 prior studies" does. Translate model output into domain terms.

03

Error Prevention

Error prevention beats error recovery. AI catching an issue before promotion is worth 10x as much as AI explaining why a broken file got downstream.

04

User Control

User control and freedom: Every AI suggestion must have a reversible action path — accept, dismiss, edit, or escalate. Never a one-way door.

05

Help & Docs

In high-stakes domains, "why did the AI flag this?" must be answerable in seconds, not by reading a model whitepaper.



The Interconnected Ecosystem



SPD

Study Protocol Designer. Utilizes AI-assisted authoring for protocol completeness, objective alignment, and strict regulatory gap analysis.



MDR

Metadata Repository. The system of record for intelligent data mapping, CDISC standards, and controlled terminology compliance workflows.



SCE

Statistical Computing Environment. Generating programs automatically from datasets & helps with TLFs for standard datasets like SDTM and ADaM prior to downstream promotion.



DSW

Data Science Workspace. Builds and publishes containerized R/Python execution environments and apps for downstream analysis. AI recommends compute sizing, diagnoses build failures, and flags package vulnerabilities to follow GxP compliance.



Takeaways

connect·share·advance

Human Centered
Clinical Data
Design



Reuse Trusted Doors AI output should walk through the same UI channels your users already understand and trust, avoiding a parallel, disconnected "AI mode."



Mandate Explainability Confidence scores and clear rationale are not nice-to-haves; in highly regulated domains, they are the difference between adoption and rejection.



AI Detects, Human Resolves Let algorithms detect early and speak up proactively, but explicitly design the UI so the human retains the final, authoritative word.



Narendra Parmar

 www.linkedin.com/in/nareniit/

 Sycamore Informatics

Sycamore Clinical Agentic Framework™

SPD  | MDR  | DTM  | DSW  | SCE  | CDR 
PROTOCOL | METADATA | TRANSFORM | OPEN-SOURCE | ANALYSIS | REPOSITORY

Sycamore Informatics

 www.sycamoreinformatics.com

connect·share·advance



**The Global Healthcare
Data Science Community**

Contact Channels

 office@phuse.global
 [phuse.global](https://twitter.com/phuse_global)

Social Media

 [/company/phuse](https://www.linkedin.com/company/phuse)
 [/phusetube](https://www.youtube.com/channel/UCphuse)
 [/phuse_global](https://twitter.com/phuse_global)
 [/phusebook](https://www.facebook.com/phusebook)



Scan & Connect !

connect·share·advance



THANK YOU!