

ZERO LAB / INITIAL CATALOG

Model cards for the work so far

*Language, memory, checked tools, integer models,
and small expert teams.*

CATALOG DATE 2026-08-23

SCOPE ILXYR AND LOCAL DEVELOP WORKSPACE

PUBLIC FLAGSHIP ZERO.4

What the scan found

This is the initial catalog from ilxylr and the local develop workspace. It groups repeated checkpoints under one model line, keeps negative results visible, and marks queued work and non-model systems so they are not mistaken for finished models.

ID	MODEL OR LINE	HOME	TYPE / SIZE	EVIDENCE STATE	CONSOLIDATION PATH
Z-001	ZERO.1	zero-grounded-literary-lm	Character MLP; 7,436 params	Retained teacher	Keep as a regression teacher and a tiny audit reference for Z5.
Z-002	ZERO.2 literary line	zero-grounded-literary-lm	Decoder-only character transformer; 4,852,992 params	Retained teacher	Keep immutable as a Z5 literary anchor and style regression control.
Z-003	ZERO Channel v1	zero-grounded-literary-lm	Channel-trained character transformer; 4.85M params	Retained experiment	Use its conversation and memory tests when freezing the Z5 baseline.
Z-004	ZERO.3	zero-grounded-literary-lm	Distilled character transformer; 4,852,992 params	Frozen teacher	Preserve as the fixed parent for ZERO.4 and a direct Z5 comparison.
Z-005	ZERO.4	zero-grounded-literary-lm	Public character transformer; 4,852,992 params; 4.9 MB int8	Public flagship	Keep as the public Z5 baseline. Do not overwrite it with an unproven candidate.
Z-006	ZERO.4 experiment branch family	zero-grounded-literary-lm	Checkpoints, adapters, and heads; 4.85M base plus small routed parts	Mixed evidence	Move checked operation heads and replay controls into the Z6 faculty track.
Z-007	Logic v1	zero-grounded-literary-lm	Logic-specialist character transformer; 4.85M params	Retained experiment	Keep as a checked Z6 faculty candidate, with the verifier remaining authoritative.
Z-008	Logic-Shakespeare v1	zero-grounded-literary-lm	Hybrid logic and literary transformer; 4.85M params	Retained experiment	Use as negative evidence for Z8 expert separation instead of one-model mixing.
Z-009	Infinite Monkey v1	zero-grounded-literary-lm	Cumulative Brainfuck, logic, and literary transformer; 4.85M params	Evidence gap	Recover the missing evidence or archive the line as an input to Z6.
Z-010	Infinite Monkey trace10m	zero-grounded-literary-lm	Systematic-execution specialist; 9,876,800 params	Evidence gap	Restore or rerun only if its result can change a Z6 design decision.
S-000	Sero 0	sero-model-lineage	Reference contract, not a trained checkpoint; 4,887,808-param reference shape	Reference contract	Keep as the stable interface behind Z5 tokenizer and data work.
S-101	Sero Latent v1	sero-model-lineage	Learned patch language model; Compute-matched small arms	No-go	Retain as a closed design branch, not as a Z5 tokenizer choice.
S-102	Sero Latent v2	sero-model-lineage	Discrete patch-code language model; Compute-matched small arms	No-go	Archive as evidence against this exact dictionary design.
S-103	Sero Latent v3	sero-model-lineage	Continuous learned-tokenizer model; 446,713 latent; 770,720 BPE control	No-go	Keep byte-BPE for Z5. Reopen latent work only after a cheaper structural repair.
S-201	Sero 1	sero-model-lineage	Dense byte-BPE base model; 6,021,312 params	Promoted baseline	Keep as the clean Z5 data and tokenizer control.
S-202	Sero 2	sero-model-lineage	Two-stage curriculum base model; 6,021,312 params	Promoted recipe	Make this the main Z5 training recipe and baseline suite.
S-203	Sero 20M	sero-model-lineage	Matched scale candidate; 20,011,136 params	Staged candidate	Treat as the first bounded Z5 scale test, not as the automatic next release.
H-001	Holo Softmax T512 control	ilxylr / holo-llm-hrr-attention	Byte GPT control; 4.330M params	Retained checkpoint	Keep as the exact attention control for future Z5 memory tests.
H-002	Holo Uniform HRR T512	ilxylr / holo-llm-hrr-attention	Uniform holographic attention GPT; 4.068M params	Retained negative result	Keep as a negative control. Do not use uniform bundling in Z5 by default.
H-003	Holo Gated HRR T512	ilxylr / holo-llm-hrr-attention	Gated holographic attention GPT; 4.069M params	Retained experiment	Use only behind explicit long-context or recall gates for Z5.

ID	MODEL OR LINE	HOME	TYPE / SIZE	EVIDENCE STATE	CONSOLIDATION PATH
H-004	Holo Gated HRR long-context line	holo-llm-hrr-attention	Long-context holographic GPT variants; About 4.07M params	Promising local evidence	Run a frozen multi-seed long-context test before any Z5 integration.
H-101	Qwen2.5-0.5B HRR conversion	holo-llm-hrr-attention	Converted external base; 0.5B class	Converted, fine-tuning	Require quality, memory, and latency gates before it can inform Z5.
H-102	Qwen2.5-1.5B HRR distillation	holo-llm-hrr-attention	Cloud conversion experiment; 1.5B class	No-go result	Archive this path until a cheaper layer-wise repair can pass a small gate.
H-103	Qwen2.5-7B HRR conversion	holo-llm-hrr-attention	Converted external base; 7B class	In progress	Pause behind the 0.5B and 1.5B preservation decision.
H-104	Qwen2.5-14B HRR conversion	holo-llm-hrr-attention	Planned conversion; 14B class	Queued, no model yet	Keep closed until smaller conversions prove value.
H-105	Qwen2.5-32B HRR QLoRA	holo-llm-hrr-attention	QLoRA conversion experiment; 32B class external base	Fine-tuning	Require an explicit stop rule and preservation baseline before more spend.
H-106	Qwen2.5-72B HRR conversion	holo-llm-hrr-attention	Planned conversion; 72B class	Queued, no model yet	Keep closed unless the full smaller-rung evidence becomes positive.
H-201	leCore HRR GPT-2	remote Hugging Face reference	External HRR reference model; Not verified locally	Remote reference	Pin and evaluate only if it can change a Holo design decision.
H-202	leCore Qwen3.5-9B Assimilated	Hugging Face, pinned by ilxvr	Assimilated external base; 9,653,104,368 params	Pinned remote model	Keep as a pinned external reference, not a Z-family parent by default.
N-001	Integer Transformer Proof v1	nsrl	Integer transformer plus suffix memory; small-h8-d128-ff256	Historical promoted proof	Retain as deployment and attribution evidence for Z6, not as the current substrate headline.
N-002	Integer Transformer Successor v2	nsrl	Transformer-only integer model; 2-layer small integer transformer	Promotion gate passed	Use as the integer substrate reference for Z6 checked tools and Z8 experts.
N-101	Literary H8 base	nsrl	Small integer literary transformer; d128, 8 heads, ff256	Promoted profile	Use as the small expert building block for Z8.
N-102	Literary H8 optimizer swarm	nsrl	Three H8 leaves plus learned routers; Three small H8 experts	Promoted experiment	Carry the expert-diversity lesson into Z8 routing tests.
N-103	Literary H8 hybrid residual experts	nsrl	Shared trunk with diagonal and rank-16 residuals; Small adapters on one H8 trunk	Promoted experts	Keep the shared-trunk expert format for Z8; redesign the router view.
N-104	Literary H8 gradient-block experts	nsrl	Gradient-clustered block experts and curriculum; Three rank-8 block experts	Promoted fixed experts	Use as the strongest local evidence for Z8 expert decomposition.
N-201	NSRLPM1 production model line	nsrl	Integer causal linear-attention model; 9,317,632 params	Generation gate failed	Do not promote. Use its exact runtime and failure audits in Z5/Z6 engineering.
N-301	NSRLTCH	nsrl	Text-conditioned bitmap denoiser; Small integer multimodal expert	Experimental	Keep as a possible narrow Z8 visual expert.
N-302	NSRLLAT1	nsrl	Prompt-to-layout latent prior; Small integer multimodal expert	Experimental	Keep only if a narrow Z8 layout task has a checked gate.
N-303	NSRLMOD1	nsrl	Discrete joint prompt/text/image-token model; Small integer multimodal model	Experimental	Treat as a sandbox input to a future Z8 multimodal expert family.
N-304	NSRLMM1	nsrl	Attention-based joint prompt/text/image model; Small native mini-transformer	Experimental target	Promote only a narrow checked expert before wider Z8 scaling.
N-401	Solomon Council vo	nsrl	Deterministic controller system, not a separate trained model; Six sealed faculties plus judge	Shadow-ready system	Use as a Z6/Z8 controller only after a fair same-model evaluation.
C-001	CRLP Prime Gap Scorer	crlp1primes	Certificate-gated neural action scorer; Tiny shared C network	Promoted evidence spine	Use the gate pattern as a Z6 checked-faculty template.

ID	MODEL OR LINE	HOME	TYPE / SIZE	EVIDENCE STATE	CONSOLIDATION PATH
C-002	CRLP Grid Path Scorer	crlprimes	Certificate-gated neural action scorer; Tiny shared C network	Promoted evidence spine	Carry the exact action-contract pattern into Z6 tools.
C-003	CRLP Delayed Claim Scorer	crlprimes	Replay-grounded neural decision scorer; Tiny shared C network	Promoted evidence spine	Use its abstention and evidence-request actions in the Z6 controller design.
C-004	CRLP multitask scorer family	crlprimes	Shared, task-head, and single-task tiny nets; 8x4 through 64x32 sweeps	Mixed research evidence	Use per-faculty evidence and selective sharing in Z6, not forced unification.
C-005	CRLP Holographic Scorer family	crlprimes	HolographicRanker, HoloFeat, Hybrid Holo NN, HoloNet, MemoryNeuron; Zero-parameter memories plus tiny neural hybrids	Sandbox family	Use Holo recall as optional Z5 memory and Z8 expert routing support, behind explicit gates.
W-001	CosyWorld Card Policy Ranker	cosyworld	Deterministic integer card ranker; 516 bytes; 24 features to 16 ReLU units	Production shadow	Use as a real product example of Z6-style model proposal plus exact execution.
X-001	holostuff LocalAgentCore	holostuff	Parameter-free vector system; No learned weights required	Supporting system	Extract only stable, licensed interfaces needed by Z5 memory and Z8 routing.
X-002	ilxvr	ilxvr	Research control plane, not a model; Protocol and evidence system	Core lab infrastructure	Make it the single evidence catalog behind Z5, Z6, and Z8.
X-003	zero-experiment-runs	zero-experiment-runs	Duplicate run storage, not a model family; Copied checkpoints and run outputs	Storage only	Deduplicate by artifact hash and link the canonical ilxvr record.
X-004	holonet	holonet	Mirror and dashboard, not a trained model; Research presentation surface	Supporting system	Link to canonical evidence instead of maintaining a second model registry.

What ZERO Lab is

ZERO Lab builds small AI systems that we can inspect, test, and run on personal hardware. ZERO is the public model family. Sero studies better language training. Holo studies memory and fixed-size attention. Solomon and NSRL study exact integer models and controllers. CRLP studies neural proposals that stay inside checked rules. ilxvr keeps the evidence record.

SMALL	Start small enough to read, rerun, and compare.
LOCAL	Keep private words on the device when the task allows it.
CHECKED	Use exact programs for facts a model should not guess.
HONEST	A no-go result closes a branch. It is not a hidden failure.

The public story should start with ZERO.4 because it is real, downloadable, and running today. The research story is larger: the lab is learning how language, memory, checked tools, integer runtimes, and expert routing fit together without pretending every experiment is a product.

Catalog labels

Public flagship means a model people can use now. Promoted baseline means a result passed its frozen research gate. Retained experiment means the model or checkpoint matters but is not the public default. No-go means the exact tested design failed. In progress and queued are not results. Supporting system means useful infrastructure, not a trained model.

One research line, four stages

The catalog suggests a clean order. Improve language first, then add checked faculties, then repeat and harden the result, then build a team of small experts.

Z5

Better local language

Start from the public ZERO.4 experience and the promoted Sero 2 data recipe. Compare the 6M control and the bounded 20M scale test. Add Holo recall only when it wins a frozen memory test.

Exit gate: Beat ZERO.4 and Sero controls on held-out language, conversation, repetition, and local-run tests across three seeds.

Z6

Language with checked faculties

Connect language to small operation heads, exact programs, CRLP-style gates, evidence requests, and abstention. Keep the controller and verifier separate from the language model.

Exit gate: A model can propose; exact tools decide; failures are visible; replay and provenance are complete.

Z7

Replication and hardening

Reserve one family number for repeats, safety and rights checks, packaging, browser and native parity, resource limits, and release rehearsal.

Exit gate: The same artifact, tests, hashes, licenses, and failure behavior reproduce from a clean checkout.

Z8

A team of small experts

Use the strongest NSRL H8 expert evidence, Holo memory, and a fair Solomon-style controller. Route only when the experts expose a real conditional gap.

Exit gate: The routed system beats the best fixed expert, preserves abstention, and reports which expert and evidence caused each decision.

First decision to freeze: the Z5 language and conversation baseline. No new Z5 training should begin until the exact data split, prompts, decoding settings, repetition measures, context-use tests, cost limits, and promotion rule are written down.

ZERO and Sero

The public line and its direct language-training successors. Historical teachers remain immutable; Sero does not rename earlier evidence.

Z-001 ZERO.1 RETAINED TEACHER	
What it is The smallest inspectable language model in the line. It acts as the frozen foundation teacher for later ZERO models. What the evidence says Frozen artifact at update 20,000. Its narrow role is foundation constraints, not broad language use.	Type Character MLP Scale 7,436 params Home zero-grounded-literary-lm Limit Very small and intentionally narrow. It is not a modern assistant. Next Keep as a regression teacher and a tiny audit reference for Z5.
EVIDENCE teachers/registry.json; README.md	
Z-002 ZERO.2 literary line RETAINED TEACHER	
What it is A literary and historical specialist trained on the project literary corpus. The frozen teacher comes from the later consolidated literary line. What the evidence says Teacher source update 12,600. The earlier literary-v6 selected checkpoint reached held-out loss 1.6641 at update 11,600.	Type Decoder-only character transformer Scale 4,852,992 params Home zero-grounded-literary-lm Limit Small specialist with a narrow source mix. It can imitate style and repeat source patterns. Next Keep immutable as a Z5 literary anchor and style regression control.
EVIDENCE teachers/registry.json; README.md, Train the literary model	
Z-003 ZERO Channel v1 RETAINED EXPERIMENT	
What it is Tests learned dialogue state, lossy memory updates, reply structure, and optional Holo recall inside a local browser runtime. What the evidence says The checkpoint and contrastive channel benchmark are retained. Holo recall is parameter-free runtime state rather than extra transformer weights.	Type Channel-trained character transformer Scale 4.85M params Home zero-grounded-literary-lm Limit The channel benchmark is narrow. This does not establish broad conversation quality. Next Use its conversation and memory tests when freezing the Z5 baseline.
EVIDENCE benchmarks/zero-channel-v1/README.md; teachers/registry.json	

Z-004 ZERO.3 FROZEN TEACHER	
What it is Combines the earlier foundation and literary functions with broader replay. It is the initialization teacher for ZERO.4. What the evidence says Frozen update 16,600. Matched validation was 1.7424, better than ZERO.2 at 1.8325 across every measured source.	Type Distilled character transformer Scale 4,852,992 params Home zero-grounded-literary-lm Limit A stronger generalist at this scale, but still narrow and weak on modern language tasks. Next Preserve as the fixed parent for ZERO.4 and a direct Z5 comparison.
EVIDENCE TEACHERS.md; teachers/registry.json	
Z-005 ZERO.4 PUBLIC FLAGSHIP	
What it is The current public ZERO model. It runs in the browser with no server and keeps the prompt on the device. What the evidence says The prospectively selected Q2.6 seed-2 update-500 artifact passed the frozen three-seed family gate. SHA-256: 44b32f2262be2754fd2eeaf16ed206bae32b4ce30d7f5541a1059cd21257ae50.	Type Public character transformer Scale 4,852,992 params; 4.9 MB int8 Home zero-grounded-literary-lm Limit Weak by modern LLM standards. External screens did not show a general language improvement over ZERO.3. The checked quantity controller is not part of the browser page. Next Keep as the public Z5 baseline. Do not overwrite it with an unproven candidate.
EVIDENCE docs/model.json; ZERO4.md; EXPERIMENTS.md	
Z-006 ZERO.4 experiment branch family MIXED EVIDENCE	
What it is Q1 through Q3.4 test quantity routing, replay protection, low-rank isolation, operation heads, and semantic transfer without hiding failed branches. What the evidence says Q2.6 produced ZERO.4. Most other branches were no-go or diagnostic. The Q3.2 operation head passed its narrow quantity package but did not replace the public model.	Type Checkpoints, adapters, and heads Scale 4.85M base plus small routed parts Home zero-grounded-literary-lm Limit Experiment checkpoints are not separate public releases. A narrow head result is not general reasoning. Next Move checked operation heads and replay controls into the Z6 faculty track.
EVIDENCE EXPERIMENTS.md; benchmarks/zero4-q*/	

Z-007 Logic v1 RETAINED EXPERIMENT	
What it is Learns verifier-generated finite proof records in the same small C transformer family. What the evidence says Selected update 26,600 with held-out next-character loss 0.1491. The small proof probe passed all six trained shapes but only one of four held-out shapes.	Type Logic-specialist character transformer Scale 4.85M params Home zero-grounded-literary-lm Limit Low character loss did not become reliable proof search or broad structural transfer. Next Keep as a checked Z6 faculty candidate, with the verifier remaining authoritative.
EVIDENCE README.md, Trained logic-v1 result	
Z-008 Logic-Shakespeare v1 RETAINED EXPERIMENT	
What it is Tests whether a logic specialist can absorb Shakespeare while keeping proof behavior. What the evidence says Selected global update 43,300 with mixed validation loss 0.8828. Dramatic language improved, while proof retention and held-out proof-shape transfer remained weak.	Type Hybrid logic and literary transformer Scale 4.85M params Home zero-grounded-literary-lm Limit Shows measurable forgetting and no reliable general proof-schema transfer. Next Use as negative evidence for Z8 expert separation instead of one-model mixing.
EVIDENCE README.md, Continued logic + Shakespeare result	
Z-009 Infinite Monkey v1 EVIDENCE GAP	
What it is Trains checked program execution, logic, and literature in cumulative stages with replay. What the evidence says The README names a completed seed-89 baseline and capacity verdict, but the referenced local result bundle is absent in the scanned checkout.	Type Cumulative Brainfuck, logic, and literary transformer Scale 4.85M params Home zero-grounded-literary-lm Limit Do not make a promotion claim until the result files and artifact hashes are restored. Next Recover the missing evidence or archive the line as an input to Z6.
EVIDENCE README.md, Train the Infinite Monkey curriculum	

Z-010 Infinite Monkey trace10m EVIDENCE GAP	
What it is Uses grouped state traces to learn primitive transitions, composed blocks, execution, synthesis, and repair. What the evidence says Training and evaluation commands are documented, but no local result bundle was found in the scanned checkout.	Type Systematic-execution specialist Scale 9,876,800 params Home zero-grounded-literary-lm Limit A documented training target is not a measured model result. Next Restore or rerun only if its result can change a Z6 design decision.
EVIDENCE README.md, systematic-execution experiment	
S-000 Sero o REFERENCE CONTRACT	
What it is Defines the lossless byte and channel-token contract that starts the Sero line without rewriting ZERO history. What the evidence says Round-trip and deterministic tokenizer checks are defined. Learned merge IDs begin at 264.	Type Reference contract, not a trained checkpoint Scale 4,887,808-param reference shape Home sero-model-lineage Limit This is an architecture and tokenizer contract, not a promoted trained model. Next Keep as the stable interface behind Z5 tokenizer and data work.
EVIDENCE docs/SERO.md; sero0-contract.json	
S-101 Sero Latent v1 NO-GO	
What it is Tests learned causal patch boundaries, a local encoder/decoder, and a global patch transformer. What the evidence says Learned boundaries improved the latent arm, but a conventional 4,096-token byte-BPE transformer won all three seeds. Preregistered comparison: 4.071 vs 4.004 BPB.	Type Learned patch language model Scale Compute-matched small arms Home sero-model-lineage Limit The run used a small, source-biased sample and charged a redundant patch-end target. Next Retain as a closed design branch, not as a Z5 tokenizer choice.
EVIDENCE docs/SERO.md; benchmarks/sero-latent-v1/RESULTS.md	

S-102 NO-GO Sero Latent v2	
What it is Tests a 4,095-entry patch dictionary plus one residual escape code against byte-BPE. What the evidence says Mean validation loss was 4.171 BPB versus 3.993 for byte-BPE across three frozen seeds.	Type Discrete patch-code language model Scale Compute-matched small arms Home sero-model-lineage Limit It tested a frequency dictionary, not a jointly learned continuous tokenizer. Rare escapes were costly. Next Archive as evidence against this exact dictionary design.
EVIDENCE docs/SERO.md; benchmarks/sero-latent-v2/RESULTS.md	
S-103 NO-GO Sero Latent v3	
What it is Runs a corrected learned-tokenizer test on balanced 100M-byte schedules with contextual chunk selection and byte reconstruction. What the evidence says At 100M bytes, mean latent quality was 2.5009 BPB versus 2.1835 for byte-BPE, 14.53% worse. Every seed failed quality.	Type Continuous learned-tokenizer model Scale 446,713 latent; 770,720 BPE control Home sero-model-lineage Limit The exact one-stage embedding-router design failed. This does not prove all learned tokenization is impossible. Next Keep byte-BPE for Z5. Reopen latent work only after a cheaper structural repair.
EVIDENCE benchmarks/sero-latent-v3/RESULTS.md	
S-201 PROMOTED BASELINE Sero 1	
What it is The first promoted dense base model in Sero. It uses the locked 4,096-entry lossless byte-BPE tokenizer. What the evidence says All three seeds passed. Mean validation BPB 1.6117 and mean test BPB 1.6181 after 135.5M token exposures per seed.	Type Dense byte-BPE base model Scale 6,021,312 params Home sero-model-lineage Limit Greedy output often loops. Promotion means reproducible pretraining, not assistant quality. Next Keep as the clean Z5 data and tokenizer control.
EVIDENCE benchmarks/sero1-pretrain-v1/RESULT.md; aggregate.json	

S-202 Sero 2 PROMOTED RECIPE	
What it is Adds broader technical, dialogue, and math data, then repairs old-source retention with a frozen consolidation stage. What the evidence says All three seeds passed every final gate. Mean final test score is 1.328497 BPB with 0.001471 sample standard deviation.	Type Two-stage curriculum base model Scale 6,021,312 params Home sero-model-lineage Limit Generation remains unreliable. The promoted object is the two-stage recipe and its evidence, not a claim of intelligence. Next Make this the main Z5 training recipe and baseline suite.
EVIDENCE benchmarks/sero2-curriculum-consolidation-replication-v1/RESULT.md	

S-203 Sero 20M STAGED CANDIDATE	
What it is Asks whether a larger model gives a clear prediction and generation gain while data, tokenizer, context, and curriculum stay fixed. What the evidence says The 6M matched baseline and frozen thresholds are present. The scanned checkout does not yet support promotion of a 20M result.	Type Matched scale candidate Scale 20,011,136 params Home sero-model-lineage Limit Scale is not a result. It must beat the frozen 6M gates and improve generation behavior. Next Treat as the first bounded Z5 scale test, not as the automatic next release.
EVIDENCE benchmarks/sero20m-scale-generation-v1/DESIGN.md; BASELINE_RESULT.md	

Holo and leCore

Memory, holographic attention, and external conversion work. Local single-seed results are kept separate from promoted evidence.

H-001 Holo Softmax T512 control RETAINED CHECKPOINT	
What it is Provides the matched softmax baseline for the Holo attention experiment on public enwik8. What the evidence says Single-seed CPU proxy: 2.9731 validation BPB and 192.99 ms per step. The checkpoint is retained locally.	Type Byte GPT control Scale 4.330M params Home ilxjr / holo-llm-hrr-attention Limit One seed on one corpus and CPU. It is a comparison control, not a public language model release. Next Keep as the exact attention control for future Z5 memory tests.
EVIDENCE ilxjr/docs/experiments/exp-006-evidence.json	

H-002 Holo Uniform HRR T512 RETAINED NEGATIVE RESULT	
What it is Replaces softmax with a fixed-size holographic state and tests constant-memory attention. What the evidence says Single-seed CPU proxy: 3.5507 BPB and 226.95 ms per step, worse and slower than softmax at T512.	Type Uniform holographic attention GPT Scale 4.068M params Home ilxjr / holo-llm-hrr-attention Limit This T512 proxy does not test the long-context scaling point where linear attention may help. Next Keep as a negative control. Do not use uniform bundling in Z5 by default.
EVIDENCE ilxjr/docs/experiments/exp-006-evidence.json	

H-003 Holo Gated HRR T512 RETAINED EXPERIMENT	
What it is Adds learned decay to bound holographic superposition and reduce crosstalk. What the evidence says Single-seed CPU proxy: 3.3308 BPB and 296.08 ms per step. It beat uniform HRR but still lost to softmax at T512.	Type Gated holographic attention GPT Scale 4.069M params Home ilxjr / holo-llm-hrr-attention Limit The decay barely moved from its initial value. Quality and speed both missed the short-context control. Next Use only behind explicit long-context or recall gates for Z5.
EVIDENCE ilxjr/docs/experiments/exp-006-evidence.json	

H-004 Holo Gated HRR long-context line		PROMISING LOCAL EVIDENCE
What it is Tests whether bounded superposition becomes useful at T2048 and under tighter decay initializations. What the evidence says At T2048, gated a=0.99 reached 3.6997 BPB versus softmax 3.9232. A separate a=0.9 sweep reached 3.2894. These are local single-seed results.	Type Long-context holographic GPT variants Scale About 4.07M params Home holo-llm-hrr-attention Limit Per-step cost was still about 1.4x softmax on CPU. The initialization acted as the real hyperparameter. Next Run a frozen multi-seed long-context test before any Z5 integration.	
EVIDENCE vendor/holo-llm-hrr-attention/RESULTS.md		
H-101 Qwen2.5-0.5B HRR conversion		CONVERTED, FINE-TUNING
What it is Swaps Qwen softmax attention for gated HRR while keeping pretrained projections. What the evidence says The repository records conversion and active fine-tuning, but no frozen promotion result was found.	Type Converted external base Scale 0.5B class Home holo-llm-hrr-attention Limit An in-progress conversion is not a released lab model and not evidence of preserved chat quality. Next Require quality, memory, and latency gates before it can inform Z5.	
EVIDENCE vendor/holo-llm-hrr-attention/README.md		
H-102 Qwen2.5-1.5B HRR distillation		NO-GO RESULT
What it is Tests a larger HRR conversion and distillation rung. What the evidence says Teacher perplexity was 8.44. The recorded final converted perplexity was 4,477.45 and the sample was badly degraded.	Type Cloud conversion experiment Scale 1.5B class Home holo-llm-hrr-attention Limit The conversion did not preserve model quality. It must not be described as a useful converted LLM. Next Archive this path until a cheaper layer-wise repair can pass a small gate.	
EVIDENCE vendor/holo-llm-hrr-attention/cloud/results_convert_Qwen2.5-1.5B.json		

H-103 Qwen2.5-7B HRR conversion IN PROGRESS	
What it is Extends the gated-HRR attention swap to a larger Qwen checkpoint. What the evidence says The local status says converting. No frozen artifact or evaluation was found.	Type Converted external base Scale 7B class Home holo-llm-hrr-attention Limit Do not spend more scale budget without a passing smaller conversion gate. Next Pause behind the 0.5B and 1.5B preservation decision.
EVIDENCE vendor/holo-llm-hrr-attention/README.md	
H-104 Qwen2.5-14B HRR conversion QUEUED, NO MODEL YET	
What it is A named rung in the HRR conversion track. What the evidence says Only queued status was found. This is not yet a developed model artifact.	Type Planned conversion Scale 14B class Home holo-llm-hrr-attention Limit Planning status must not be counted as a result. Next Keep closed until smaller conversions prove value.
EVIDENCE vendor/holo-llm-hrr-attention/README.md	
H-105 Qwen2.5-32B HRR QLoRA FINE-TUNING	
What it is Tests parameter-efficient recovery after the attention conversion at a large external scale. What the evidence says The repository records active QLoRA fine-tuning. No frozen quality result was found.	Type QLoRA conversion experiment Scale 32B class external base Home holo-llm-hrr-attention Limit High cost and weak smaller-rung evidence make this a research branch, not a lab flagship. Next Require an explicit stop rule and preservation baseline before more spend.
EVIDENCE vendor/holo-llm-hrr-attention/README.md	

H-106 Qwen2.5-72B HRR conversion QUEUED, NO MODEL YET	
What it is The largest named rung in the conversion track. What the evidence says Only queued status was found. No local model artifact exists in the scan. EVIDENCE vendor/holo-llm-hrr-attention/README.md	Type Planned conversion Scale 72B class Home holo-llm-hrr-attention Limit No authorization should come from the catalog alone. Next Keep closed unless the full smaller-rung evidence becomes positive.
H-201 leCore HRR GPT-2 REMOTE REFERENCE	
What it is Serves as prior art and a from-scratch proof point for the Holo conversion work. What the evidence says A remote model link is present, but no local artifact or lab evaluation was found. EVIDENCE vendor/holo-llm-hrr-attention/README.md	Type External HRR reference model Scale Not verified locally Home remote Hugging Face reference Limit Reference material is not local promotion evidence. Next Pin and evaluate only if it can change a Holo design decision.
H-202 leCore Qwen3.5-9B Assimilated PINNED REMOTE MODEL	
What it is A public Qwen3.5-derived checkpoint registered by exact Hub revision and file hashes. What the evidence says ilxjr records revision 4e14e0ee3d5b6936dfd3ddofa7454d9118fe88c5 and about 19 GB of weights. Import proves identity only. EVIDENCE ilxjr/examples/schema/huggingface-model.json; docs/HUGGINGFACE.md	Type Assimilated external base Scale 9,653,104,368 params Home Hugging Face, pinned by ilxjr Limit The local control plane did not download or evaluate the weights. Quality claims remain external. Next Keep as a pinned external reference, not a Z-family parent by default.

NSRL and Solomon

Integer models, small literary experts, multimodal components, and bounded controllers. Exact replay is strong; language quality remains the main open gate.

N-001 Integer Transformer Proof v1 HISTORICAL PROMOTED PROOF	
What it is Proves deterministic integer training, replay, and packaging for an assisted transformer system. What the evidence says 579 per-mille accuracy on 5,896 targets. Later ablation showed suffix memory caused all top-1 gain, while transformer logits improved probability error.	Type Integer transformer plus suffix memory Scale small-h8-d128-ff256 Home nsrl Limit This is not a transformer-only win. It must be named as an assisted memory system. Next Retain as deployment and attribution evidence for Z6, not as the current substrate headline.
EVIDENCE benchmarks/integer-transformer-proof-v1/promoted-candidate.json; component-ablation.json	
N-002 Integer Transformer Successor v2 PROMOTION GATE PASSED	
What it is Repairs the v1 attribution gap with a physically unassisted transformer trained on canonical integer base-2 NLL. What the evidence says 25,347,655 total NLL millibits on 5,896 targets, better than uniform, retrieval, byte n-gram, and the matched float transformer. Zero zero-probability windows.	Type Transformer-only integer model Scale 2-layer small integer transformer Home nsrl Limit Top-1 mistakes remain high. This is a substrate result, not a language-quality release. Next Use as the integer substrate reference for Z6 checked tools and Z8 experts.
EVIDENCE benchmarks/integer-transformer-proof-v1/successor-v2-evidence.json	
N-101 Literary H8 base PROMOTED PROFILE	
What it is Tests whether more small heads improve the fixed-width integer literary model without making it much larger. What the evidence says On the matched 512-window probe, H8 reached 189 per-mille accuracy and 53,753 mean error versus H2 at 68 and 61,343. Replay was byte-identical.	Type Small integer literary transformer Scale d128, 8 heads, ff256 Home nsrl Limit Next-byte metrics improved, but generated prose still failed quality gates. Next Use as the small expert building block for Z8.
EVIDENCE nsrl/docs/literary-scale-up.md	

N-102 Literary H8 optimizer swarm PROMOTED EXPERIMENT	
What it is Creates useful expert diversity through different integer optimizer scales, then routes by target-blind context. What the evidence says The token router raised accuracy from 148 to 151 per mille and avoided 102 mistakes. A 16-token router made a smaller gain with fewer switches. EVIDENCE data/experiments/literary-h8-swarm-v1/report.json; docs/literary-scale-up.md	Type Three H8 leaves plus learned routers Scale Three small H8 experts Home nsrl Limit The gains are narrow and do not prove prose quality. Next Carry the expert-diversity lesson into Z8 routing tests.
N-103 Literary H8 hybrid residual experts PROMOTED EXPERTS	
What it is Reduces repeated trunk work while keeping several bounded expert corrections. What the evidence says The hybrid fixed expert reached 57,818 mean error, 557 Q15 better than the earlier diagonal fixed result. The router collapsed to almost all Blake and was not promoted. EVIDENCE nsrl/docs/literary-scale-up.md; recursive-literary-swarm-experiment.md	Type Shared trunk with diagonal and rank-16 residuals Scale Small adapters on one H8 trunk Home nsrl Limit Expert gain did not become successful target-blind routing or good generation. Next Keep the shared-trunk expert format for Z8; redesign the router view.
N-104 Literary H8 gradient-block experts PROMOTED FIXED EXPERTS	
What it is Builds expert groups from frozen-trunk gradient signatures instead of author labels, then continues them through bounded stages. What the evidence says All three fixed experts beat the trunk. After curriculum, the best fixed expert improved untouched-final error by 5,506 Q15; token-oracle headroom grew to 18,360. EVIDENCE data/experiments/literary-h8-gradient-block-*/report.json; docs/literary-scale-up.md	Type Gradient-clustered block experts and curriculum Scale Three rank-8 block experts Home nsrl Limit Learned routers still missed the best fixed expert and prose gates still failed. Next Use as the strongest local evidence for Z8 expert decomposition.

N-201 NSRLPM1 production model line GENERATION GATE FAILED	
What it is The production-scale integer language-model path with exact cache, resumable optimization, and extensive numeric-health audits. What the evidence says Serving and integrity gates pass, but open generation fails. The latest documented row has severe one-token loops, weak diversity, and no measured context use.	Type Integer causal linear-attention model Scale 9,317,632 params Home nsrl Limit Good replay and numeric health do not imply language quality. Next Do not promote. Use its exact runtime and failure audits in Z5/Z6 engineering.
EVIDENCE nsrl/PROJECT_STATUS.md; docs/production-model-v1.md	
N-301 NSRLTCH EXPERIMENTAL	
What it is Maps text conditions and noisy bitmap state toward a deterministic u8 ink rendering surface. What the evidence says The architecture and training path exist inside the Solomon multimodal suite. It is experiment evidence, not a headline model.	Type Text-conditioned bitmap denoiser Scale Small integer multimodal expert Home nsrl Limit No broad image-quality or product promotion claim is supported by the catalog scan. Next Keep as a possible narrow Z8 visual expert.
EVIDENCE nsrl/README.md; docs/schemas.md	
N-302 NSRLLAT1 EXPERIMENTAL	
What it is Learns a prompt and layout prior for the Solomon visual pipeline. What the evidence says The training format and artifact contract exist. No separate public promotion was found.	Type Prompt-to-layout latent prior Scale Small integer multimodal expert Home nsrl Limit A component artifact is not a complete text-to-image model. Next Keep only if a narrow Z8 layout task has a checked gate.
EVIDENCE nsrl/README.md; docs/schemas.md	

N-303 NSRLMOD1		EXPERIMENTAL
What it is Models one serialized stream of prompt, text, and coarse image tokens. What the evidence says The repo includes training and sampling support, but no independent public promotion claim was found.	Type Discrete joint prompt/text/image-token model Scale Small integer multimodal model Home nsrl Limit Coarse joint-token experiments do not establish high-quality multimodal generation. Next Treat as a sandbox input to a future Z8 multimodal expert family.	
EVIDENCE nsrl/README.md; docs/schemas.md		

N-304 NSRLLMM1		EXPERIMENTAL TARGET
What it is A causal attention wrapper for joint prompt, text, and image tokens and the named target for a narrow NSRL-born expert. What the evidence says Architecture, training, browser sampling, and artifact checks are present. Project status still calls promotion a future gate.	Type Attention-based joint prompt/text/image model Scale Small native mini-transformer Home nsrl Limit It is not yet the promoted product model. Next Promote only a narrow checked expert before wider Z8 scaling.	
EVIDENCE nsrl/README.md; PROJECT_STATUS.md		

N-401 Solomon Council vo		SHADOW-READY SYSTEM
What it is Coordinates model faculties inside explicit permission, evidence, tool, and resource circles. It can select, ask, request evidence, or abstain. What the evidence says Core receipts replay exactly. A later hardening audit found the original solo/council comparison lacked tool parity, so council promotion is not authorized.	Type Deterministic controller system, not a separate trained model Scale Six sealed faculties plus judge Home nsrl Limit Council behavior must not be credited as a new foundation model. Next Use as a Z6/Z8 controller only after a fair same-model evaluation.	
EVIDENCE nsrl/PROJECT_STATUS.md; benchmarks/solomon-council-v1/hardening-result.json		

CRLP and CosyWorld

Small action scorers that work inside exact or replayable rules, including one production shadow model.

C-001 CRLP Prime Gap Scorer PROMOTED EVIDENCE SPINE	
What it is Ranks candidate prime gaps after algebraically certified modular rules remove impossible actions. What the evidence says Exact next-prime verification owns correctness. The learned scorer ranks the residual candidate set under replayable traces.	Type Certificate-gated neural action scorer Scale Tiny shared C network Home crlprimes Limit The model proposes and ranks; the exact verifier decides. Next Use the gate pattern as a Z6 checked-faculty template.
EVIDENCE crlprimes/README.md; results/latest/summary.md	

C-002 CRLP Grid Path Scorer PROMOTED EVIDENCE SPINE	
What it is Ranks four legal moves after certified topology and bound rules protect optimal actions. What the evidence says BFS owns the exact shortest-path successor label. The scorer operates only inside the checked action surface.	Type Certificate-gated neural action scorer Scale Tiny shared C network Home crlprimes Limit Grid success is a narrow exact task, not general planning. Next Carry the exact action-contract pattern into Z6 tools.
EVIDENCE crlprimes/README.md; results/latest/summary.md	

C-003 CRLP Delayed Claim Scorer PROMOTED EVIDENCE SPINE	
What it is Ranks trust, ignore, request, and cautious actions under a deterministic hidden-outcome simulator and finite safety checks. What the evidence says The production direct gate resolves clean trust/request bands. Stress rows expose noise sensitivity and the limit of indirect proxies.	Type Replay-grounded neural decision scorer Scale Tiny shared C network Home crlprimes Limit This is deterministic replay authority, not proof of real-world trustworthiness. Next Use its abstention and evidence-request actions in the Z6 controller design.
EVIDENCE crlprimes/README.md; results/latest/summary.md	

C-004 CRLP multitask scorer family MIXED RESEARCH EVIDENCE	
What it is Tests transfer, adapter heads, capacity, leave-one-out behavior, and gate invariance across promoted and sandbox domains. What the evidence says The evidence shows some positive transfer, but no single shared pattern wins uniformly across prime gap, grid path, and delayed claim.	Type Shared, task-head, and single-task tiny nets Scale 8x4 through 64x32 sweeps Home crlprimes Limit Do not turn mixed multitask results into a universal shared-model claim. Next Use per-faculty evidence and selective sharing in Z6, not forced unification.
EVIDENCE results/2026-05-17-multitask-roadmap/	

C-005 CRLP Holographic Scorer family SANDBOX FAMILY	
What it is Tests content-addressable action memory, structured receptive fields, hybrid neural scoring, and capacity-partitioned memories. What the evidence says Local-continuity tasks can reach 100% with holographic k-NN, while discontinuous delayed claims favor the neural scorer. Hybrid results are task-dependent.	Type HolographicRanker, HoloFeat, Hybrid Holo NN, HoloNet, MemoryNeuron Scale Zero-parameter memories plus tiny neural hybrids Home crlprimes Limit Most variants remain sandbox mechanisms. They need verifier-backed rows and matched baselines before promotion. Next Use Holo recall as optional Z5 memory and Z8 expert routing support, behind explicit gates.
EVIDENCE docs/holonet-retrospective.md ; docs/gap-analysis.md	

W-001 CosyWorld Card Policy Ranker PRODUCTION SHADOW	
What it is Scores every legal resident card with shared weights and keeps world authority in the deterministic game kernel. What the evidence says The fixed-seed 10,000-avatar top-3 regression found treasure in every episode. The model is deployed in shadow mode with exact hash and schema checks.	Type Deterministic integer card ranker Scale 516 bytes; 24 features to 16 ReLU units Home cosyworld Limit Synthetic bootstrap success is not a general gameplay policy result. Live authority remains bounded. Next Use as a real product example of Z6-style model proposal plus exact execution.
EVIDENCE cosyworld/models/card-policy/README.md ; v2/docs/card-policy-ranker.md	

Supporting systems and duplicate stores

Items found during the scan that matter to the catalog but are not separate trained models.

X-001 holostuff LocalAgentCore SUPPORTING SYSTEM	
What it is Provides deterministic hypervectors, binding, storage, and recall used as reference material by Holo and channel-memory work. What the evidence says The scan found a broad vector and holographic toolkit, not one promoted trained model. EVIDENCE holostuff/REFERENCE.md; holographic/	Type Parameter-free vector system Scale No learned weights required Home holostuff Limit Do not list the whole toolkit as an LLM or claim learned model quality. Next Extract only stable, licensed interfaces needed by Z5 memory and Z8 routing.
X-002 ilxyr CORE LAB INFRASTRUCTURE	
What it is Freezes experiments, admits runs, records evidence, imports pinned model identities, and keeps negative results visible. What the evidence says It carries the promoted experiment lane and the first recovered Holo comparison, plus ZERO replication records. EVIDENCE ilxyr/README.md; docs/PROGRAM.md	Type Research control plane, not a model Scale Protocol and evidence system Home ilxyr Limit A control plane does not create model quality by itself. Next Make it the single evidence catalog behind Z5, Z6, and Z8.
X-003 zero-experiment-runs STORAGE ONLY	
What it is Keeps execution outputs from ZERO branches. What the evidence says The scan found duplicate candidate checkpoints, not a distinct trained model line. EVIDENCE zero-experiment-runs/q28-*/	Type Duplicate run storage, not a model family Scale Copied checkpoints and run outputs Home zero-experiment-runs Limit Copies must not become new catalog entries or new promotion claims. Next Deduplicate by artifact hash and link the canonical ilxyr record.

X-004
holonet

SUPPORTING SYSTEM

What it is
Presents holographic research notes and experiments.

What the evidence says
No separate promoted weight artifact was found in this repository.

Type Mirror and dashboard, not a trained model
Scale Research presentation surface
Home holonet

Limit
The presentation layer must not duplicate CRLP or Holo model identities.

Next
Link to canonical evidence instead of maintaining a second model registry.

EVIDENCE holonet/quantum-neural-holography-report.md

Where the catalog came from

The report is an evidence-backed local scan, not a claim that every worktree is clean or every remote artifact is public. These repository revisions and files supplied the main facts.

REPOSITORY	SCANNED REVISION	KEY EVIDENCE
ilxvr	e92382ff2a5e8714466533a160f6609b4ef9cee8	README.md; docs/experiments/exp-006-evidence.json; docs/HUGGINGFACE.md
sero-model-lineage	88a7ad5104a3719f7d99f7d9c1f2cec93863b055	docs/SERO.md; TEACHERS.md; Sero result bundles; ZERO experiment ledger
nsrl	94515eabc301e759226095a14fada7020d1c5dd8	PROJECT_STATUS.md; docs/integer-transformer-proof-v1.md; docs/literary-scale-up.md
crllprimes	cd50ed3e90bea7947ef71e340aed74a6c729c82a	README.md; results/latest/; docs/holonet-retrospective.md
cosyworld	962d9ea439a0c7c7cdb6cfe57c4399f397b1300c	models/card-policy/README.md; v2/docs/card-policy-ranker.md
holostuff	2639825e49b69326af0abbff20cf3b254644e58f	REFERENCE.md; holographic/
zero-grounded-literary-lm local clone	dc72668a781a2e742c0349f833b98a4af56714c8	Historical local clone; public Pages release prepared from the current upstream clone

Scope notes

Repeated worktrees, copied ZERO run directories, and hundreds of NSRL experimental checkpoints are grouped under their model line. Third-party provider models named inside product configuration are not counted as lab-developed models. Queued conversion rungs are listed so the plan is visible, but they are marked as not yet models. Remote Hugging Face artifacts are separated from locally retained weights.

What to do next

Turn this PDF into the first machine-readable catalog. Give every entry one stable ID, one owning repository, one parent, one artifact hash or explicit no-artifact state, one evidence state, one public claim boundary, and one Z5/Z6/Z8 decision it can change. ilxvr should hold that registry while each project keeps its own training code.