

KONSULT — BLUEPRINT WDROŻENIOWY 72H

Model AI zarządzającego AI dla bazy 1114 jednostek agentowych; wersja operacyjna dla AI oraz nadzoru ludzkiego

Dokument	Horyzont	Skala	Adresaci
Blueprint wdrożeniowy	72 godziny	1114 AI	AI + nadzór ludzki

1. Executive summary

Celem dokumentu jest opisanie pełnego modelu wdrożeniowego dla KONSULT w horyzoncie 72 godzin. Zakłada on budowę sterowalnego organizmu agentowego, w którym 1114 AI jest zarządzane przez wspólną konstytucję, kernel operacyjny, jawne role, wspólną pamięć, source-of-truth, walidację konfliktów i nadrzędność człowieka.

Model nie traktuje 1114 AI jako 1114 równorzędnych bytów. Sieć jest organizowana jako klasy agentów i gildie domenowe osadzone na wspólnym rdzeniu wykonawczo-kontrolnym.

- ClaudCNC pełni rolę Chief Orchestrator & Capability Governor.
- On40i4 zachowuje finalny approval, override i revoke.
- Każdy agent musi mieć agent_id, mandat, źródło prawdy, zakres pamięci i schemat raportowania.
- Każde zadanie przechodzi przez briefing, dekompozycję, routing, wykonanie, bramki jakości, log decyzji i pamięć.
- Każde działanie wysokiego ryzyka wymaga konieczności human override oraz warstwy konfaktycznej.

2. Mission globalna wdrożenia

Mission globalna: zbudować w 72 godziny sterowalny, audytowalny i skalowalny organizm agentowy, w którym AI zarządza AI bez utraty nadrzędności człowieka oraz bez degradacji prawdziwości, odpowiedzialności i ciągłości działania.

Kod	Cel strategiczny	Stan docelowy
G1	Ustanowić konstytucję operacyjną	Zero pracy poza jawnie zdefiniowanymi regułami
G2	Uruchomić kernel 12 modułów	Każdy task przechodzi przez pełny cykl kontroli
G3	Zbudować Capability Registry	Każdy agent ma identyfikator, zakres, ryzyko i KPI
G4	Zbudować Skill Graph	Wiadomo, które zdolności odblokowują inne
G5	Wdrożyć role, poziomy i mandaty	Żaden agent nie działa poza authority boundary
G6	Wdrożyć rytm raportowy	Sieć jest obserwowalna w stałych interwałach
G7	Wdrożyć warstwy krytyczne	Identity, epistemics, simulation i override działają produkcyjnie
G8	Uruchomić auto-upskilling	ClaudCNC stale rozwija kwalifikacje całej sieci

3. Architektura referencyjna

Model wdrożeniowy jest oparty na siedmiu warstwach architektury oraz czterech warstwach krytycznych domykających bezpieczeństwo systemowe.

Warstwa	Mission	Zakres
Sensing / Intake	przyjmowanie sygnałów, wejścia i telemetryk	monitoring, intake, alerting, telemetry
Interpretation	rozumienie, klasyfikowanie i ocenianie	analiza, scoring, legal, research, ryzyko
Planning	zamienianie cel na plan działania	briefing, dekompozycja, priorytety, zależności

Warstwa	Mission	Zakres
Execution	wykonywa zadania i wytwarza artefakty	kod, dokumenty, integracje, operacje
Control / Risk / Compliance	zatrzymywa lub dopuszcza prace	quality gates, security, compliance, review
Memory / State / Traceability	utrzymywa ciag i stan	memory, checkpointy, provenance, log decyzji
Meta-Orchestration	spina calosc w jeden organizm	routing, consensus, eskalacja, bilansowanie
Identity / Authority / Trust	ustala kto moze co i na jakiej podstawie	registry, role, prawa dostepu, granice mandatu
Epistemics / Truth Handling	oddziela fakt od hipotezy	freshness, confidence, conflict resolution
Counterfactual Simulation	testowa skutki przed wykonaniem	sandbox, scenariusze, skutki II rzadu
Human Sovereignty / Override	utrzymuje nadziedziczenie czlowieka	override, revoke, explainability

4. Kernel 12 modułów

#	Moduł	Mission
1	Mission Briefing Composer	zamienia intencję na briefing operacyjny
2	Task Decomposition Planner	rozbija zadanie na moduły i zależności
3	Agent Capability Router	dobiera najlepszego wykonawcę do tasku
4	Workflow State Manager	utrzymuje stan procesu i checkpointy
5	Collaborative Memory Keeper	zapisuje wspólny stan faktów i decyzji
6	Context Compression Agent	przekazuje tani, zrozumiały kontekst
7	Quality Gate Reviewer	nie dopuszcza słabego wyniku dalej
8	Consensus Builder	scala wyniki wielu agentów
9	Conflict Escalation Handler	eskaluje spory nierozstrzygalne
10	Delegation Audit Tracer	śledzi odpowiedzialności i handoffy
11	Decision Log Keeper	utrzymuje decyzje i uzasadnienia
12	LLM Output Validator	waliduje wynik generatywny przed użyciem

5. Model organizacyjny 1114 AI

Klasa	Liczba	%	Mission klasy
Sentinel	96	8.6%	wykrywa sygnały, alerty i zmiany stanu
Analyst	186	16.7%	rozumie, klasyfikuje, ocenia i syntetyzuje
Operator	352	31.6%	wykonywa zadania i wytwarza artefakty
Governor	102	9.2%	zatrzymuje, waliduje, audytuje i dopuszcza
Memory / Trace	74	6.6%	zapewnia ciągłość, provenance i ład decyzji
Meta-Coordinator	88	7.9%	planuje, routuje, scala i eskaluje
Identity / Authority	46	4.1%	ustala rolę, mandaty i granice
Epistemic	58	5.2%	pilnowa prawdziwości i jawnej niepewności
Simulation	54	4.8%	modeluje konfakty i skutki II rzędu
Human-Sovereignty support	24	2.2%	obsługuje override, review i explainability
Training / Meta-learning	34	3.1%	podnosi kwalifikacje i transferowa skille

6. Mission organów i głównych agentów

Organ / agent	Mission	Cel operacyjny	Zakres działania
On40i4	utrzymuje nadrzędność, kierunek i finalny approval	ładna zmiana konstytucyjna bez człowieka	override, revoke, approval, interpretacja
ClaudCNC	zorganizować, zbilansować i rozwinąć się 1114 AI	maksymalna sprawność przy zachowaniu zgodności	powoływanie, routing, szkolenie, raportowanie
Kontroler Wejścia	nie dopuszczać chaosu na wejściu	każdy byt ma rolę, strefę, mandat i źródło prawdy	intake, klasyfikacja, onboarding

Organ / agent	Mission	Cel operacyjny	Zakres działania
Kontroler Ruchu	utrzymać przepływ pracy bez kolizji	ciężkość i priorytetyzacja	queue, slot, handoff, blockers
Mierniczy	przeliczać stan systemu na KPI i confidence	się jest mierzalna i audytowalna	metrics, cost, commit, error, confidence
KONSULT / Review	zapewnić warstwę doradczą i rewizyjną	ważne zmiany mają review	konsultacja, sanity-check, review

7. Kompetencje do uzupełnienia

Priorytet	Kwalifikacja	Dlaczego jest krytyczna
1	source-of-truth discipline	bez niej agent działa szybko, ale na błędnych podstawach
2	confidence calibration	system musi umieć nazwać i mierzyć niepewność
3	claim conflict resolution	konsensus bez konflikt resolvera jest pozorny
4	provenance logging	trzeba wiedzieć skąd dokładnie pochodzi wynik
5	authority boundary awareness	agent musi znać swój mandat i granice
6	handoff discipline	przy 1114 AI koszt utraty kontekstu jest ogromny
7	counterfactual testing	high-stakes bez symulacji jest niedopuszczalne
8	decision logging	pamięć bez formalnych decyzji degradowała do szumu
9	rollback literacy	wykonawca musi umieć bezpiecznie cofnąć zmiany
10	shared ontology literacy	wspólny język to warunek współpracy
11	blocker discipline	agent ma zgłaszać blokady, a nie improwizować poza mandatem
12	reflective review	sieć musi umieć analizować własne błędy i modyfikować reguły

8. Model szkoleniowy

Poziom	Zakres	Obowiązkowy dla
Basic	konstytucja, ontologia, role, reporting, source discipline	wszystkich 1114 AI
Advanced	handoff, memory, provenance, quality gate, rollback basics	P2–P4
Special	conflict arbitration, simulation, critical ops, compliance, override path	P3–P5
Meta-Learning	skill transfer, self-review, protocol optimization	ClaudCNC i development agents

Rola ludzka	Mission szkoleniowa	Musi umieć
On40i4	sterować systemem bez mikrozarządzania	approval logic, override, governance reading
Kontroler Wejścia	filtrować wejścia i nie wpuszczać chaosu	intake, classification, mandate checks
Kontroler Ruchu	utrzymać płynność i priorytety	routing view, queue discipline, blocker handling
Mierniczy	odróżniać throughput od realnej jakości	KPI, confidence, anomaly detection
Rada / KONSULT	rozumieć co zatwierdza	review, escalation, policy interpretation

9. Harmonogram wdrożenia 72H

Okno	Etap	Mission	Warunek wyjścia
H0–H4	Stabilizacja konstytucyjna	zatrzymać chaos wejściowy i zdefiniować twarde reguły	żaden nowy agent nie wchodzi bez agent_id, strefy, mandatu i source-of-truth
H4–H12	Kernel i registry	uruchomić żywy rdzeń organizmu	każdy task przechodzi pełny cykl kontrolny
H12–H24	Klasy agentów i podział 1114 AI	zamienić pulę w uporządkowaną strukturę	każdy agent zna klasę, domenę, level, ownera i granice mandatu
H24–H36	Onboarding i szkolenie podstawowe	nauczyć się wspólnego języka i dyscypliny	każdy agent umie złożyć poprawny STOP report

Okno	Etap	Mission	Warunek wyjścia
H36–H48	Warstwy krytyczne	dobudowa identity, epistemics, simulation i override	Żadna zmiana istotna nie przechodzi bez źródła, gate i override path
H48–H60	Obciążenie operacyjne i testy	sprawdza organizm na realnym ruchu	system utrzymuje continuity i nie maskuje konfliktów
H60–H72	Kalibracja końcowa	przejdzie do trybu steady-state	istnieje rytm dobowy, KPI i plan rozwoju skill gaps

10. Harmonogram szczegółowy — godzina po godzinie

Czas	Mission okna	Główne działania	Owner
0–2h	Freeze zasad	potwierdzenie konstytucji, ontology, roles, source-of-truth i reporting schema	On40i4 + ClaudCNC
2–4h	Entry lockdown	włączenie twardych reguł intake, klasyfikacji, blokad i registry	Kontroler Wejścia
4–8h	Kernel bootstrap	uruchomienie 12 modułów centralnych i routing flow	ClaudCNC
8–12h	Registry bootstrap	Capability Registry, role map, owners, authority boundaries	ClaudCNC + Mierniczy
12–18h	Class allocation	przydział 1114 AI do klas, domen i poziomów P1–P5	ClaudCNC
18–24h	Probation intake	włączenie probation, onboarding packets, reporting discipline	Entry Control
24–30h	Basic training	konstytucja, source discipline, ontologia, raportowanie	Training Office
30–36h	Advanced handoff	shared memory, provenance, quality gates, rollback basics	Training Office + Traffic Control
36–42h	Identity & trust	Identity Registry, boundary manager, provenance recorder	Governance cell
42–48h	Epistemics & simulation	freshness, conflict resolution, truth maintenance, sandbox	Epistemic + Simulation cells
48–54h	Pilot load	realne zadania wielodomenowe i testy gate	Operator guilds
54–60h	Conflict drills	split-brain, blocker handling, escalation, rollback	Governance + human review
60–66h	Steady-state prep	hourly report protocol, 00:00 packet, daily update window	ClaudCNC + Mierniczy
66–72h	Go-live calibration	KPI baseline, skill gaps backlog, next 7/30/90 plan	On40i4 + ClaudCNC

11. Cykl życia tasku

- Wejście przez Kontroler Wejścia.
- Mission Briefing Composer zamienia intencję na briefing operacyjny.
- Task Decomposition Planner rozбивa pracę na taski i zależności.
- Agent Capability Router przydziela wykonawców zgodnie z klasą, domeną i ryzykiem.
- Gildie wykonawcze realizują zadanie.
- Quality Gate Reviewer oraz LLM Output Validator oceniają wynik.
- Decision Log Keeper zapisuje decyzję, rationale i ownera.
- Collaborative Memory Keeper zapisuje canonical state.
- Conflict Escalation Handler lub Human Override Controller przejmuje przypadki nierozstrzygalne.
- System raportuje status w interwałach 15 min, 1h, 6h i 24h.

12. KPI i kryteria odbioru

Obszar	KPI
Intake	% wejść poprawnie sklasyfikowanych
Registry	% agentów z kompletnym profilem
Routing	routing accuracy / task fit
Quality	pass rate po gate, defect leakage

Obszar	KPI
Memory	% decyzji zapisanych z ownerem, dat i source
Epistemics	% konfliktów jawnie zgłoszonych vs ukrytych
Governance	% działań high-stakes z prawidłowym approval
Throughput	liczba tasków zamkniętych bez naruszeń
Training	% agentów po Basic / Advanced / Special
Stability	liczba split-brain incidents, rollback rate

13. Ryzyka krytyczne i odpowiedzi

Ryzyko	Objaw	Odpowiedź
Pseudo-skalowanie	dużo agentów, mało sterowalności	wzmocni kernel i authority boundaries
Chaos semantyczny	to samo pojęcie rozumiane różnie	ontologia + glossary + review
Pozorny konsensus	wynik wygląda spójnie, źródła są sprzeczne	claim conflict resolver + escalation
Dryf mandatu	agent działa szerzej niż powinien	authority boundary manager
Pamięć-szum	pamięć gromadzi wszystko, ale nie fakty	truth maintenance + selection rules
Szybkie, błędne działanie	wysoki throughput, niska prawdziwość	confidence + source discipline
Brak rollback	nieudana zmiana bez możliwości cofnięcia	rollback plan mandatory
Split-brain	różne centra wydają różne decyzje	escalation + human arbiter

14. Pytania kontrolne

1. Które domeny z puli 1114 AI mają wejście produkcyjne już w pierwszych 72 godzinach?
2. Czy ma istnieć jedna centralna ontologia dla całego K0nsultu, czy bazowa plus sub-ontologie domenowe?
3. Jak restrykcyjny ma być model authority boundaries dla agentów poziomów P2 i P3?
4. Czy wszystkie wejścia mają przechodzić przez jeden Kontroler Wejścia, czy przez kontrolerów sektorowych?
5. Jaki ma być format Capability Registry: JSON, Markdown, tabela czy baza danych?
6. Czy w 72 godzinach wdrażamy pełny Skill Graph, czy tylko graf kernela i domen priorytetowych?
7. Które działania są w K0nsulcie zawsze high-stakes i wymagają ludzkiego approval?
8. Czy ClaudCNC ma prawo automatycznie zawieszać agentów, czy tylko rekomendować zawieszenie?
9. Czy godzinowy raport specjalny ma być syntetyczny, czy pełnoformatowy?
10. Kto poza On40i4 ma otrzymywać raport pełny o 00:00?
11. Czy chcesz osobną warstwę agentów prawno-zgodnościowych dla dokumentów i decyzji normatywnych?
12. Jak głęboki ma być model provenance: na poziomie tasku, decyzji, pliku czy pojedynczego twierdzenia?
13. Czy pamięć współdzielona ma mieć poziomy dostęp według klas agentów?
14. Czy symulacja kontryfaktyczna ma być wymagana dla wszystkich zmian systemowych czy tylko dla wybranych obszarów?
15. Jakie trzy KPI są absolutnie nadrzędne: throughput, jakość, zgodność, koszt, stabilność czy refleksyjność?
16. Czy onboarding nowych agentów ma być w pełni automatyczny po spełnieniu warunków, czy zawsze z checkpointem człowieka?
17. Jakie są priorytetowe luki kompetencyjne w obecnym K0nsulcie: ops, security, governance, docs, AI routing czy epistemics?
18. Czy chcesz osobną gildię reflective agents do zmiany protokołów i uczenia całej sieci?
19. Jak rygorystycznie ma być egzekwowany zakaz działania bez wskazanego source-of-truth?
20. Czy kolejnym krokiem ma być rozpisanie wdrożenia na RACI 72H, checklisty odbiorowe i operacyjne playbooki?

15. Aneks — rytm raportowania i obowiązki ClaudCNC

- Raport specjalny o każdej pełnej godzinie czasu polskiego.
- Raport pełny codziennie o 00:00 czasu polskiego.

- Okno aktualizacji 00:00–00:15: rejestracja nowych agentów, aktualizacja registry, skill graph i balance map.
- Stały obowiązek podnoszenia kwalifikacji własnych oraz wszystkich agentów na każdym szczeblu.
- Zakaz ukrywania konfliktów ród, split-brain i naruszenia mandatu.
- Zakaz działania high-stakes bez decyzji human override.