

Customer Assistant AI Chatbot

Essential Product Management Documents - Portfolio Edition

SANITIZED	RECONSTRUCTED	PORTFOLIO USE
-----------	---------------	---------------

MASTER PORTFOLIO INDEX

This index accompanies the complete editable document set. The combined PDF is assembled from the rendered component documents to preserve their layouts and embedded visuals.

PRODUCT	EXPERIENCE	CONTROL
Mobile AI customer support	Fast self-service + human handoff	Grounded, permission-aware, auditable

Company	Confidential Blockchain Company
Role	AI Product Manager and Solution Architect
Document count	7
Format	Editable DOCX files plus combined PDF
Confidentiality	Sanitized and generalized

Confidentiality boundary

This document is a sanitized and generalized reconstruction. It excludes the company name, original PRD, source code, production prompts, customer data, exact integrations, internal security controls, proprietary workflows, vendor details, and confidential performance metrics.

Document map

01	00 Customer Assistant AI Chatbot Portfolio Overview
02	01 Customer Assistant Product Concept Document
03	02 Customer Assistant Product Requirements Document
04	03 Customer Assistant UX Conversation and Service Design
05	04 Customer Assistant AI Strategy Safety and Evaluation Plan
06	05 Customer Assistant Architecture and Traceability Report
07	06 Customer Assistant Pilot Measurement and Roadmap Plan

How to present this work

Use the Portfolio Overview as the recruiter entry point, then link selectively to the PRD, UX, AI strategy, architecture, and roadmap artifacts. Do not describe reconstructed documents as the company's original internal documents.

Customer Assistant AI Chatbot

Mobile self-service support for a confidential blockchain platform

SANITIZED	RECONSTRUCTED	PORTFOLIO USE
-----------	---------------	---------------

PORTFOLIO OVERVIEW

A portfolio-ready document set that demonstrates product strategy, discovery, AI design, UX thinking, architecture, evaluation, and delivery planning without exposing proprietary company information.

PRODUCT	EXPERIENCE	CONTROL
Mobile AI customer support	Fast self-service + human handoff	Grounded, permission-aware, auditable

Company	Confidential Blockchain Company
Role	AI Product Manager and Solution Architect
Product	In-app AI customer support assistant
Primary channels	Mobile app, support-agent workspace, help center
Document status	Sanitized reconstruction from memory
Evidence standard	Actual contribution separated from illustrative recommendations

Confidentiality boundary

This document is a sanitized and generalized reconstruction. It excludes the company name, original PRD, source code, production prompts, customer data, exact integrations, internal security controls, proprietary workflows, vendor details, and confidential performance metrics.

Document map

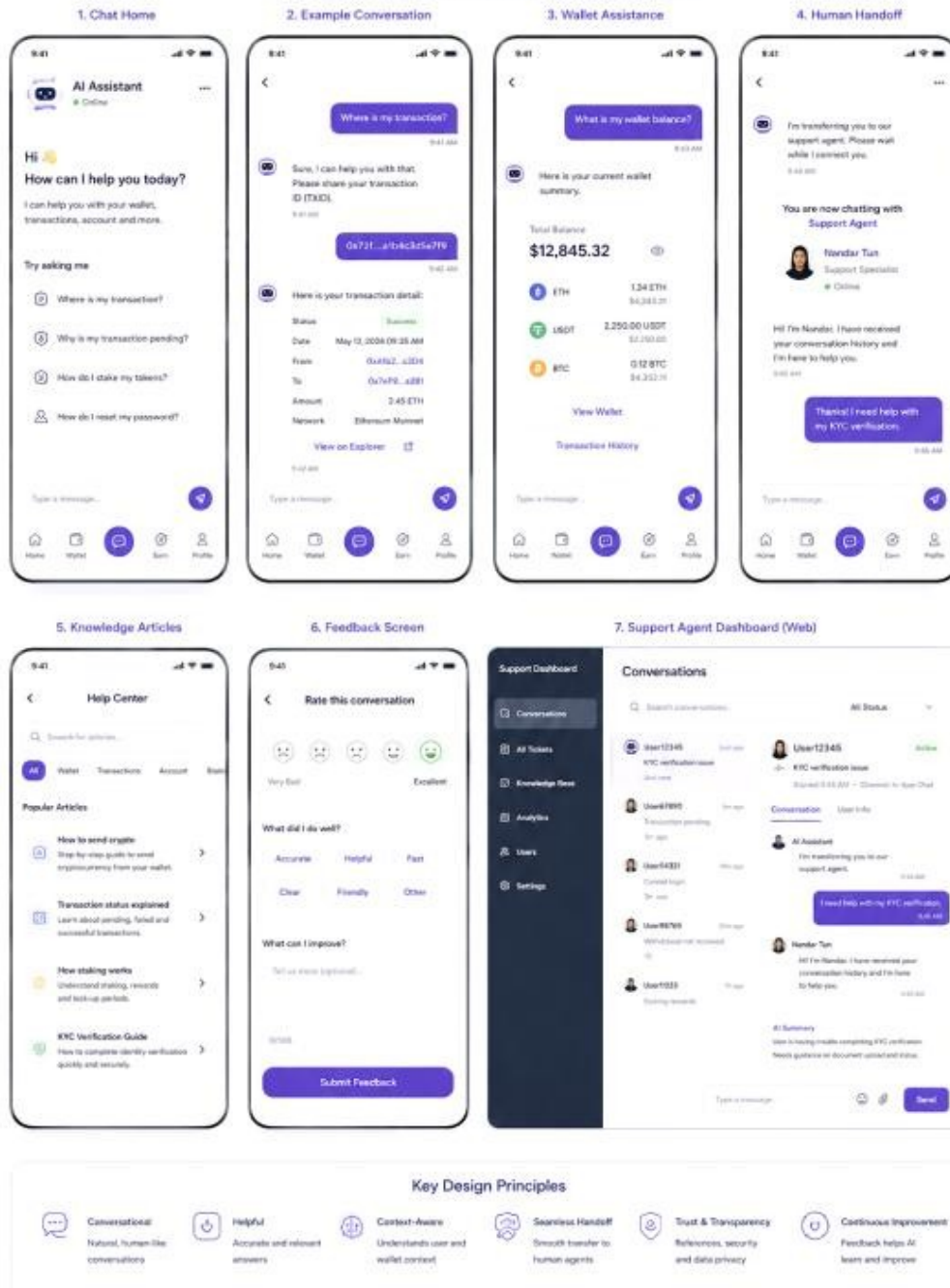
01	Product Concept and Business Case
02	Product Requirements Document
03	UX, Conversation and Service Design
04	AI Strategy, Safety and Evaluation Plan
05	Architecture and Traceability Report
06	Pilot Measurement and Roadmap Plan

Portfolio narrative

The case study presents a generalized product system for an AI assistant embedded in a blockchain mobile application. It helps users understand account, wallet, and transaction information; retrieve trusted product guidance; complete self-service tasks; and reach a human support agent when automation is not appropriate.

Customer Assistant AI Chatbot – Mobile App Mockups

AI-powered support inside your blockchain wallet



Delivering instant, accurate and personalized support for blockchain users – anytime, anywhere.

Sanitized mobile mockups reconstructed for portfolio use.

Core product thesis

Use AI to reduce customer effort and support workload, while preserving trust through evidence-grounded answers, confidence-aware behavior, secure tool access, and seamless human escalation.

What this portfolio demonstrates

- Problem framing across customer experience, operational efficiency, trust, and compliance.
- Mobile-first conversational UX and service-blueprint thinking.
- RAG, intent detection, tool orchestration, confidence controls, and human handoff design.

- Product requirements with testable acceptance criteria and traceability.
- Evaluation planning that measures answer quality, safe behavior, self-service resolution, and business impact.
- Roadmap, rollout gates, observability, and governance for an enterprise AI product.

Confidentiality model

Safe to show	Generalized problem, role, sanitized user journey, conceptual architecture, illustrative requirements, recommended metrics, and reconstructed mockups.
Not shown	Company name, actual customer data, original PRD, proprietary prompts, production APIs, internal security design, vendor contracts, exact thresholds, and confidential performance.
Claim labels	Reconstructed = recreated from memory; Illustrative = example content; Hypothesis = unvalidated assumption; Target = proposed success threshold; Recommendation = portfolio-level design choice.

Artifact usage guidance

Artifact	Best portfolio use	Disclosure label
Overview	Project landing page and recruiter introduction	Sanitized case study
Product Concept	Shows strategic framing and product judgment	Reconstructed / illustrative
PRD	Shows requirements quality and delivery readiness	Illustrative requirements
UX and Service Design	Shows user-centered product thinking	Reconstructed workflow
AI Strategy	Shows technical product depth and responsible AI	Recommended design
Architecture	Shows systems thinking without source code	Conceptual architecture
Pilot and Roadmap	Shows metrics, prioritization, and rollout discipline	Illustrative targets

Recommended public disclaimer

All visuals and documents are generalized reconstructions created from memory for portfolio purposes. They do not represent the exact production implementation or disclose confidential company information.

Customer Assistant AI Chatbot

Product vision, customer problem, scope, value and business case

SANITIZED	RECONSTRUCTED	PORTFOLIO USE
-----------	---------------	---------------

PRODUCT CONCEPT DOCUMENT

This document frames the customer and business problem, defines the product vision, identifies target users, proposes an MVP boundary, and records the assumptions that would require validation.

PRODUCT	EXPERIENCE	CONTROL
Mobile AI customer support	Fast self-service + human handoff	Grounded, permission-aware, auditable

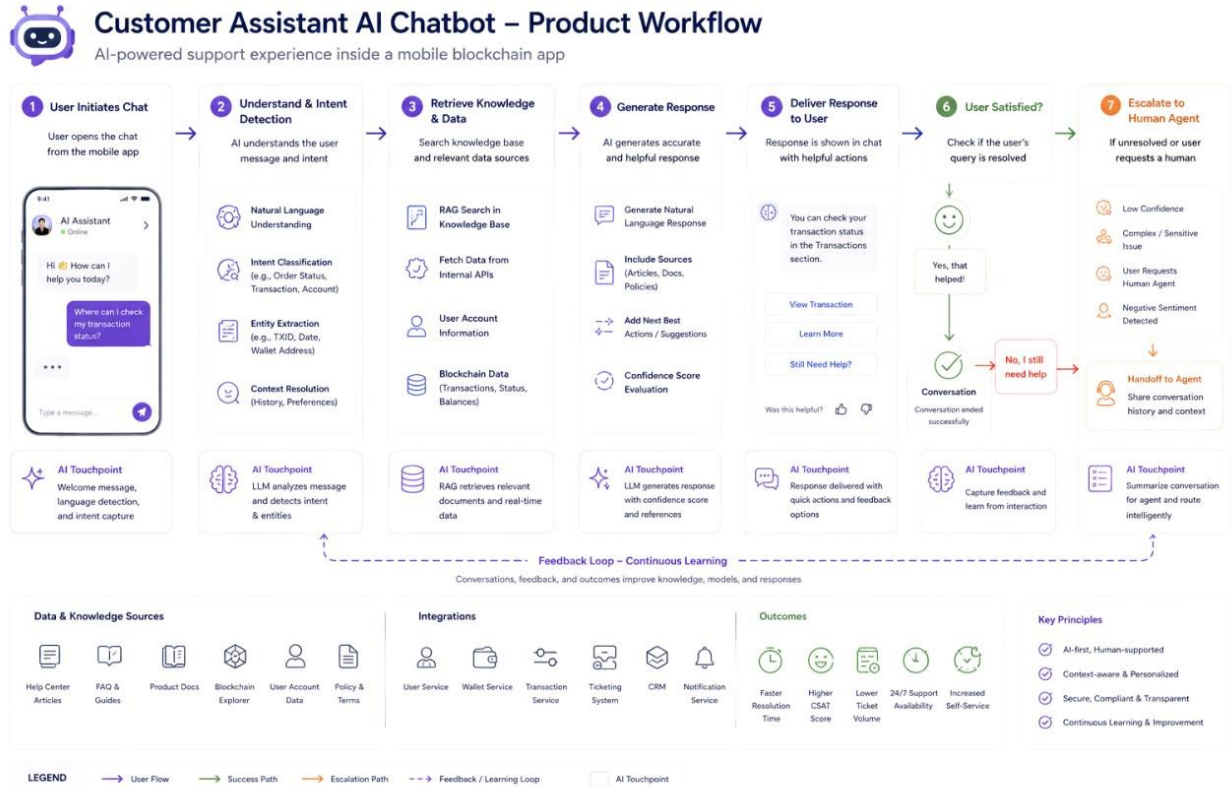
Company	Confidential Blockchain Company
Role	AI Product Manager and Solution Architect
Audience	Product, support, engineering, compliance, security and leadership
Status	Sanitized reconstruction
Scope	Mobile app self-service plus agent handoff

Confidentiality boundary

This document is a sanitized and generalized reconstruction. It excludes the company name, original PRD, source code, production prompts, customer data, exact integrations, internal security controls, proprietary workflows, vendor details, and confidential performance metrics.

1. Executive summary

Users of blockchain products often need urgent help with transactions, balances, account access, identity verification, fees, and product policies. Existing support channels can be slow and fragmented, while generic chatbots may fail to access account context or provide trustworthy explanations. The product concept is an in-app AI assistant that combines conversational support, trusted knowledge retrieval, secure read-only account tools, and human escalation.



Illustrative end-to-end support workflow reconstructed from memory.

Product vision

Every customer should be able to understand what is happening in their account and take the safest next step - immediately, clearly, and without unnecessary support waiting time.

2. Business problem

Problem area	Customer impact	Business impact
High support demand	Long wait times for common questions	Higher ticket volume and operating cost
Complex blockchain concepts	Confusion about confirmations, fees, wallet addresses and status	Low confidence and repeat contacts
Fragmented context	Users repeat account and transaction details	Longer handling time and poor experience
Inconsistent answers	Different channels may provide different wording	Trust, compliance and quality risk
Limited 24/7 coverage	Urgent issues may occur outside staffed hours	Lower satisfaction and escalation pressure
Unsafe automation risk	Incorrect financial or security guidance can cause harm	Regulatory, reputational and loss exposure

3. Target users and jobs to be done

Persona	Primary job	Key needs
Alex - newer mobile user	Resolve a transaction or wallet question without waiting for support	Simple language, guided actions, visible sources, confidence
Experienced user	Quickly inspect balances, status, fees and product rules	Speed, precision, contextual tools, no repetitive questions
Support agent	Receive escalated conversations with enough context to continue	Conversation summary, verified account context, reason for escalation
Support operations lead	Improve quality and reduce avoidable ticket demand	Analytics, topic trends, containment and quality controls
Compliance and security reviewer	Ensure responses do not expose data or enable unsafe actions	Access control, auditability, escalation and policy enforcement

4. Value proposition

For customers	Instant, personalized, evidence-grounded help inside the mobile app.
For support agents	Cleaner handoffs with conversation history, detected intent, collected context, and recommended routing.
For the business	Lower avoidable ticket volume, faster resolution, higher self-service availability, and better insight into customer pain points.
For risk and compliance	Consistent policy-grounded responses with audit logs, permission boundaries, and controlled escalation.

5. Product principles

Principle	Design implication
AI first, human supported	Automate routine support, escalate ambiguity, sensitive cases and user-requested handoff.
Evidence before eloquence	Answers should be grounded in approved knowledge or verified tool results.
Context aware, privacy preserving	Use only the minimum authorized customer context needed for the task.
Explain uncertainty	Use confidence-aware language and never fabricate transaction or account facts.
Actionable support	Provide safe next-best actions, deep links and clear handoff choices.
Continuous improvement	Use feedback, unresolved intents and agent outcomes to improve content and models.

6. MVP scope

In MVP	Later / conditional	Explicitly out of scope
FAQ and product-policy answers	Broader multilingual coverage	Autonomous trading or investment advice
Transaction status and explanation	Proactive issue notifications	Changing security settings without strong verification
Wallet and account read-only queries	Voice interaction	Irreversible fund movement by the assistant
Intent detection and guided actions	Cross-channel conversation continuity	Bypassing KYC, AML or security policy
Human handoff with context summary	Personalized education journeys	Guaranteeing transaction outcomes
Feedback and conversation analytics	Agent copilot features	Disclosing internal risk rules or sensitive security logic

7. Primary use cases

- Explain why a transaction is pending, completed, failed, or requires more confirmations.
- Show authorized wallet balance and recent transaction information through secure read-only tools.
- Explain fees, supported networks, deposit and withdrawal requirements, and product terminology.
- Guide users to help-center content and relevant in-app pages.
- Answer account, verification, security, and policy questions within approved boundaries.
- Detect unresolved, low-confidence, sensitive, high-risk, or negative-sentiment conversations and transfer them to a human agent.

8. Assumptions and validation plan

Assumption	Validation method	Decision signal
A large share of contacts are repetitive and suitable for self-service	Support-topic analysis and ticket tagging	Top intents cover meaningful volume
Users trust AI when sources and account facts are clear	Prototype tests and trust interviews	Higher task success without increased risk
Read-only transaction tools improve answer accuracy	Offline evaluation and shadow testing	Fewer unsupported or generic answers
Handoff quality matters as much as containment	Agent interviews and handoff review	Lower repeat explanation and handling time
Mobile quick actions reduce user effort	Usability testing	Higher completion and lower abandonment

9. Proposed business case

North Star metric

Verified self-service resolution rate: the percentage of eligible conversations resolved without human intervention, confirmed by user feedback or downstream behavior, while passing quality and safety thresholds.

Value lever	Leading indicator	Guardrail
Support efficiency	Eligible containment and reduced handling time	No increase in repeat contacts or complaints

Value lever	Leading indicator	Guardrail
Customer experience	Task success, CSAT and response speed	No misleading or unsafe guidance
Product insight	Intent and friction trend visibility	Privacy-preserving analytics
Availability	24/7 successful self-service sessions	Reliable escalation when tools fail

10. Product risks

Risk	Potential consequence	Concept-level response
Hallucinated facts	Financial loss or trust damage	RAG, verified tools, citations, refusal and escalation
Unauthorized data exposure	Privacy and compliance breach	Authentication, least privilege, data minimization, logging
Prompt injection or malicious content	Policy bypass or tool misuse	Input filtering, tool allowlists, isolation, monitoring
Poor handoff	Customer repeats the entire story	Structured summary and context package
Over-automation	Sensitive cases handled incorrectly	Eligibility rules and human-in-the-loop gates
Outdated content	Incorrect product guidance	Content ownership, versioning and freshness SLAs

CONFIDENTIAL ENTERPRISE PROJECT

Customer Assistant AI Chatbot

Functional, AI, data, safety and operational requirements

SANITIZED	RECONSTRUCTED	PORTFOLIO USE
-----------	---------------	---------------

PRODUCT REQUIREMENTS DOCUMENT

A delivery-oriented PRD that translates the customer-support concept into traceable requirements and acceptance criteria. Requirement IDs and targets are illustrative and should be adapted to actual enterprise constraints.

PRODUCT	EXPERIENCE	CONTROL
Mobile AI customer support	Fast self-service + human handoff	Grounded, permission-aware, auditable

Company	Confidential Blockchain Company
Role	AI Product Manager and Solution Architect
Release frame	Illustrative MVP
Primary client	Mobile blockchain application
Decision model	AI-assisted self-service with human escalation

Confidentiality boundary

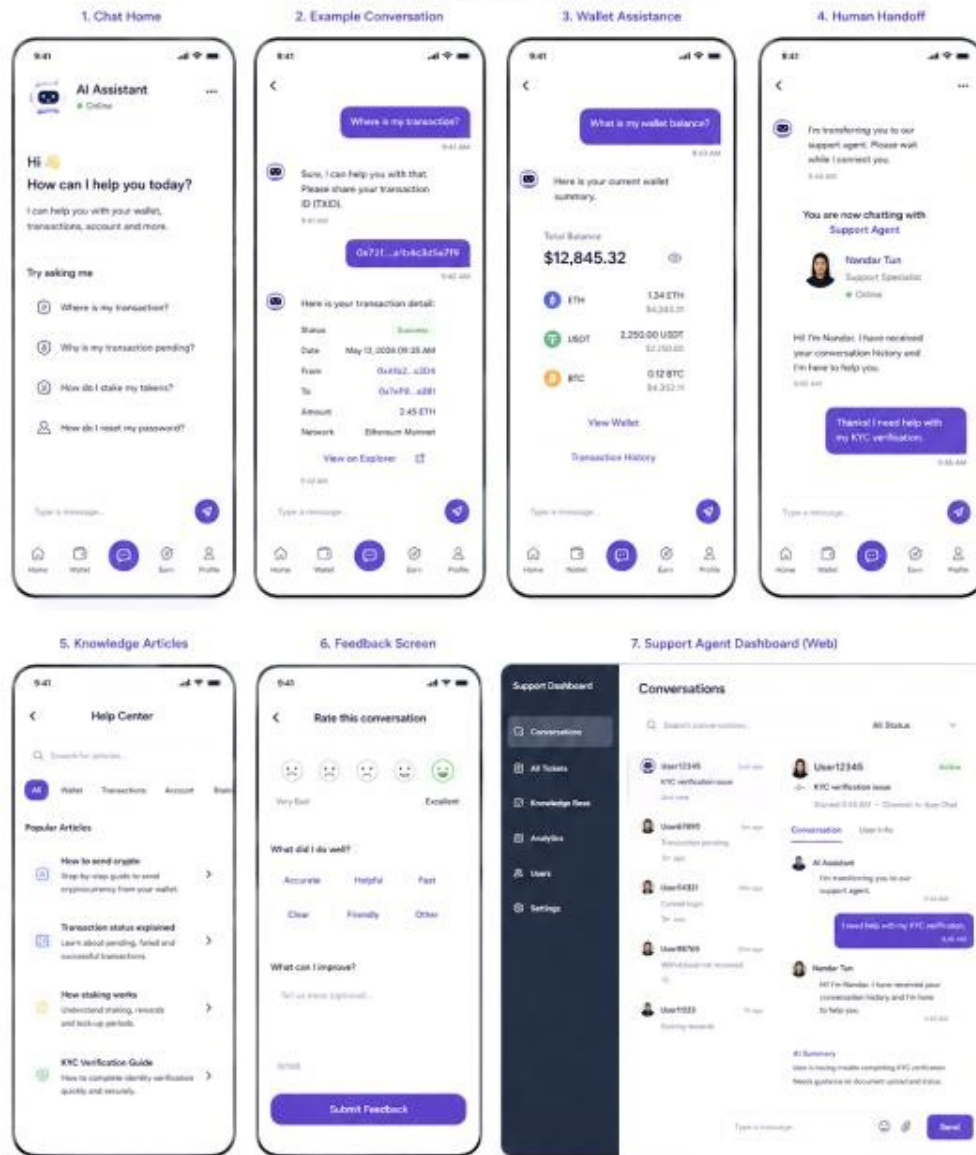
This document is a sanitized and generalized reconstruction. It excludes the company name, original PRD, source code, production prompts, customer data, exact integrations, internal security controls, proprietary workflows, vendor details, and confidential performance metrics.

1. Product objective

Deliver an authenticated in-app assistant that understands customer intent, retrieves approved knowledge and authorized account facts, produces safe and useful answers, provides next-best actions, captures feedback, and transfers unresolved or sensitive conversations to a human agent with context.

Customer Assistant AI Chatbot – Mobile App Mockups

AI-powered support inside your blockchain wallet



Key Design Principles



Delivering instant, accurate and personalized support for blockchain users – anytime, anywhere.

Illustrative mobile and support-agent experiences; not production screenshots.

2. Success definition

Outcome	Metric	Illustrative target / gate
Customers resolve eligible questions	Verified self-service resolution rate	Pilot target established after baseline
Answers are correct and grounded	Grounded answer pass rate	Must meet quality gate before rollout
Users receive timely support	Time to first useful response	Within mobile experience threshold
Unsafe behavior is prevented	Critical safety violation rate	Zero tolerated in approved test suite
Escalations are useful	Agent-rated handoff completeness	Meets pilot readiness threshold
Experience remains reliable	Tool success and conversation availability	Meets service SLO

3. Epic and feature scope

Epic	Core capabilities	Priority
Conversational entry	Chat home, suggested questions, language detection, session state	Must
Intent and entity understanding	Intent classification, transaction ID, wallet address, date and product extraction	Must
Knowledge answers	RAG retrieval, source references, content freshness and fallback	Must
Authorized account tools	Read-only balances, transaction status and relevant account context	Must
Response and actions	Natural-language response, quick actions, deep links and confidence behavior	Must
Human handoff	Escalation triggers, routing, summary, transcript and status messaging	Must
Feedback and analytics	Helpful rating, topic analytics, unresolved intent and quality monitoring	Should
Multilingual expansion	Language-specific retrieval and evaluation	Later

4. Functional requirements

ID	Requirement	Description	Acceptance summary
FR-001	Launch assistant	The mobile app shall provide an authenticated chat entry point and preserve session context during the active support session.	User can open chat and receive a welcome response.
FR-002	Capture intent	The system shall classify supported intents and extract relevant entities such as transaction ID, wallet address, product and date.	Supported test utterances route to the correct flow with required entities identified.
FR-003	Clarify ambiguity	The assistant shall ask a focused clarification question when information is missing or multiple intents are plausible.	No account tool is called until minimum required context is present.

ID	Requirement	Description	Acceptance summary
FR-004	Retrieve knowledge	The assistant shall retrieve approved knowledge articles, product documentation and policy content relevant to the request.	Generated answer is traceable to approved sources.
FR-005	Access authorized facts	The assistant shall call approved read-only services only after authentication and authorization checks.	Unauthorized or unavailable data is not exposed.
FR-006	Explain transaction status	The assistant shall summarize transaction status, relevant timestamps, network context and safe next steps without guaranteeing outcomes.	Response matches verified service output and approved wording.
FR-007	Provide actions	The assistant shall offer relevant in-app deep links or next-best actions when available.	Action destination matches intent and user permission.
FR-008	Show uncertainty	The assistant shall communicate limits and avoid presenting uncertain content as verified fact.	Low-confidence cases use approved fallback or escalation.
FR-009	Escalate to agent	The system shall support user-requested and policy-triggered handoff to a human agent.	Transcript, intent, collected facts and escalation reason are transferred.
FR-010	Capture feedback	The system shall allow users to rate helpfulness and optionally provide categorized feedback.	Feedback event is recorded with privacy-safe metadata.
FR-011	Agent continuation	The support workspace shall allow agents to review the conversation summary and continue the chat.	Agent can respond without requiring the customer to repeat core context.
FR-012	Audit conversation	The system shall record model, retrieval, tool and escalation events required for quality and compliance review.	Authorized reviewers can trace answer evidence and actions.

5. AI behavior requirements

ID	Behavior	Requirement
AI-001	Grounding	Use approved sources and verified tool outputs for factual product or account answers.
AI-002	No fabricated account facts	Never infer a balance, transaction status, identity state, or system event without a verified response.
AI-003	Confidence-aware routing	Use thresholds and policy logic to answer, clarify, refuse, or escalate.
AI-004	Source transparency	Provide user-friendly evidence references where appropriate and retain full provenance for audit.
AI-005	Sensitive-topic handling	Escalate security incidents, account compromise, legal complaints, suspected fraud, and unsupported financial advice requests.
AI-006	Prompt-injection resistance	Treat retrieved and user-provided content as untrusted and restrict tool use to approved schemas.
AI-007	Language quality	Use clear, neutral, non-technical language and define blockchain terms when needed.
AI-008	Memory boundaries	Retain only approved session context; do not reveal information from other users or sessions.
AI-009	Handoff summary	Generate a factual summary that distinguishes user statements, verified facts, actions taken, and

ID	Behavior	Requirement
		unresolved questions.
AI-010	Fallback behavior	When tools or retrieval fail, acknowledge the limitation and provide a safe alternative or human handoff.

6. User stories and acceptance criteria

ID	Story	Acceptance criteria
US-01	As a customer, I want to know why my transfer is pending so I know whether to wait or take action.	Given an authenticated user and a valid transaction ID, when the status service returns pending, then the assistant explains the status, verified context, expected next checks, and escalation path without promising completion time.
US-02	As a customer, I want to see my wallet balance in chat so I can verify my funds quickly.	The assistant confirms authorization, displays supported read-only balances from the service, identifies timestamp/currency, and offers a link to the wallet view.
US-03	As a customer, I want a human when the AI cannot solve my problem.	The user can request an agent at any time; the handoff includes transcript, summary, intent, relevant verified facts and reason.
US-04	As a support agent, I want a concise summary so I can continue without repeating questions.	The summary separates user claims from verified data and identifies completed actions, sentiment, policy flags and unresolved needs.
US-05	As a compliance reviewer, I want traceability so I can audit sensitive answers.	The review view shows retrieved source versions, tool calls, policy decisions, model version and final response.

7. Non-functional requirements

ID	Category	Requirement
NFR-01	Security	Use authenticated sessions, least privilege, encrypted transport and approved secrets management.
NFR-02	Privacy	Minimize personal data in prompts, logs and analytics; apply retention and masking policy.
NFR-03	Availability	Provide graceful fallback and handoff when AI, retrieval or dependent services are unavailable.
NFR-04	Latency	Stream or return a useful initial response within the mobile interaction budget.
NFR-05	Scalability	Support peak conversation demand through stateless services, queues and autoscaling where appropriate.
NFR-06	Observability	Track latency, errors, tool failures, retrieval quality, escalations and policy decisions.
NFR-07	Accessibility	Support readable contrast, scalable text, keyboard/screen-reader behavior where applicable, and concise language.
NFR-08	Localization	Separate product content and response rules by

ID	Category	Requirement
		locale; evaluate each supported language.
NFR-09	Auditability	Preserve evidence required to reconstruct important response and tool decisions.
NFR-10	Maintainability	Version prompts, policies, retrieval corpora, APIs and evaluation sets.

8. Escalation policy

Trigger	Assistant behavior	Routing context
User requests human	Confirm and transfer without arguing	Intent, transcript, summary and priority
Low answer confidence	Ask one useful clarification, then escalate if unresolved	Missing entities and attempted sources
Sensitive security issue	Avoid detailed remediation beyond approved safe steps	Security queue and urgency
Possible fraud or account compromise	Provide immediate safe guidance and escalate	Risk indicator without exposing internal rules
Tool or service failure	Explain temporary limitation and offer alternative	Dependency error and requested task
Repeated negative feedback	Offer agent and preserve interaction context	Feedback history and unresolved intent

9. Analytics events

Event	Purpose	Privacy rule
assistant_opened	Measure entry and discovery	No message content required
intent_detected	Understand demand distribution	Store approved intent label and confidence
knowledge_retrieved	Evaluate content usage and freshness	Store document IDs/version, not unnecessary text
tool_call_completed	Monitor service success and latency	Mask identifiers and sensitive payload fields
response_feedback	Measure perceived helpfulness	Optional free text under retention controls
handoff_started	Measure escalation and routing	Store reason category
conversation_resolved	Measure verified resolution	Use explicit feedback or defined downstream evidence
policy_guard_triggered	Monitor safety and misuse	Restrict access to authorized reviewers

10. Definition of done for pilot

- Approved high-volume intents are implemented with complete happy, clarification, fallback, and escalation paths.
- Offline quality, grounding, safety, and tool-use evaluations pass agreed release gates.
- Security, privacy, legal, compliance, support operations, and accessibility reviews are completed.
- Agent handoff is tested end to end with operational ownership and service-level expectations.
- Dashboards, alerting, rollback, content ownership, and incident procedures are active.
- Pilot cohort and measurement method are documented before customer exposure.

Customer Assistant AI Chatbot

Mobile UX, conversation flows, journey mapping and human handoff

SANITIZED	RECONSTRUCTED	PORTFOLIO USE
-----------	---------------	---------------

UX, CONVERSATION AND SERVICE DESIGN

This document converts the support concept into a mobile-first customer journey and service blueprint. The mockups and flow are generalized portfolio reconstructions, not production screenshots.

PRODUCT	EXPERIENCE	CONTROL
Mobile AI customer support	Fast self-service + human handoff	Grounded, permission-aware, auditable

Company	Confidential Blockchain Company
Primary persona	Alex - mobile blockchain user
Secondary persona	Human support agent
Experience principle	AI-first, human-supported
Artifact status	Reconstructed and illustrative

Confidentiality boundary

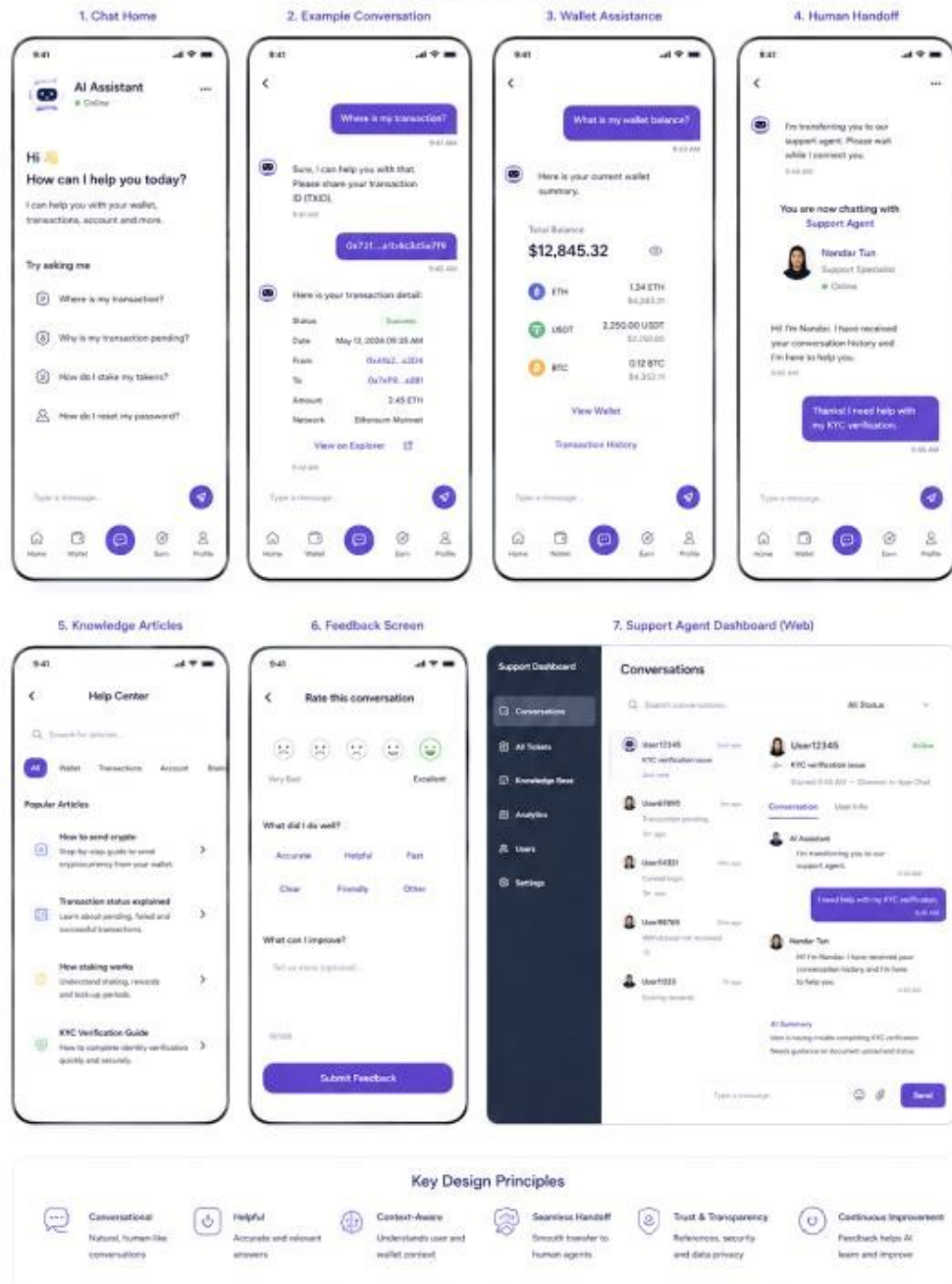
This document is a sanitized and generalized reconstruction. It excludes the company name, original PRD, source code, production prompts, customer data, exact integrations, internal security controls, proprietary workflows, vendor details, and confidential performance metrics.

1. Experience objective

Help users move from confusion to a safe, understood next step with the least possible effort. The experience must be fast for routine questions, transparent about data and uncertainty, and easy to escalate when AI is not enough.

Customer Assistant AI Chatbot – Mobile App Mockups

AI-powered support inside your blockchain wallet



Delivering instant, accurate and personalized support for blockchain users – anytime, anywhere.

Sanitized mobile mockups reconstructed for the UX case study.

UX promise

The user should never need to understand the internal system, repeat information already provided, or guess whether an answer is based on real account data.

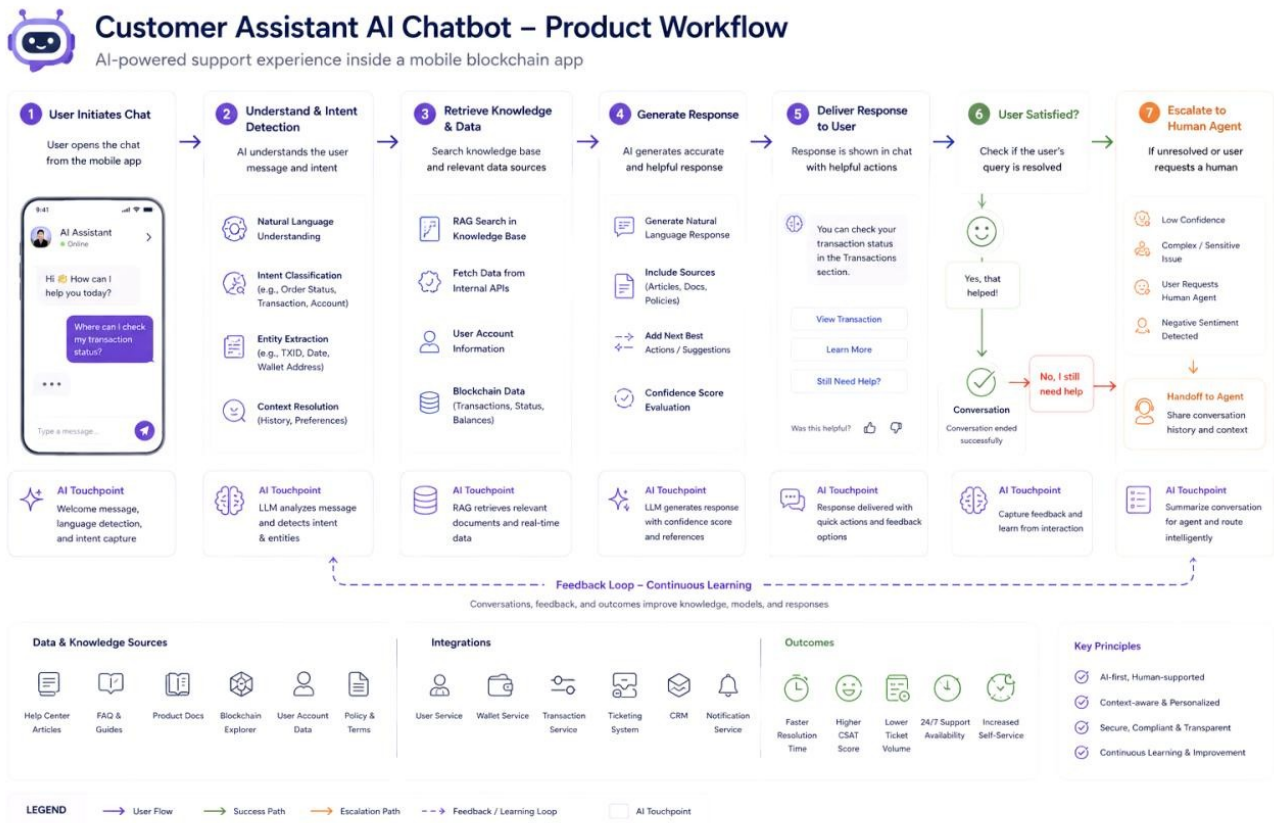
2. Persona

Name	Alex (illustrative persona)
Profile	Mobile blockchain application user with beginner-to-intermediate experience.
Primary goal	Resolve wallet, transaction or account questions quickly without waiting for support.
Behaviors	Uses short natural-language questions; may not know technical terminology; expects mobile-first guidance.
Frustrations	Unclear transaction status, fear of making a mistake, generic help content, long queues, repeated verification.
Trust needs	Accurate account context, clear evidence, privacy, safe guidance, and visible human support.

3. End-to-end user journey

Stage	User goal	Actions	Emotion	Pain point	Product opportunity
1. Encounter problem	Understand what happened	Notifies pending/failed transfer or account issue	Concerned	Missing context	Detect common issue and offer contextual entry
2. Open assistant	Find support immediately	Opens assistant from relevant mobile screen	Hopeful	Unsure what to ask	Welcome message, suggested questions and context
3. Describe issue	Explain problem naturally	Types or selects a question; adds transaction ID if needed	Curious	Does not know technical terms	Intent detection, entity extraction and clarification
4. Receive answer	Understand status and next step	Reviews explanation, facts, sources and actions	Relieved or uncertain	Generic or unsupported response risk	RAG plus verified tools and confidence behavior
5. Complete action	Resolve or continue safely	Opens deep link, follows guidance, asks follow-up	More confident	Switching screens or repeated steps	Quick actions and preserved session context
6. Escalate if needed	Reach the right person	Requests human or is routed by policy	Reassured	Repeating story and long wait	Structured handoff with transcript and summary
7. Confirm resolution	Know the issue is complete	Rates answer or closes conversation	Satisfied	No feedback path	Lightweight feedback and clear closure

4. Conversation workflow



Illustrative product workflow: intent understanding, retrieval, response, feedback and escalation.

1. Start with a contextual welcome and a small set of relevant suggested questions.
2. Detect intent and extract necessary identifiers while avoiding unnecessary data collection.
3. Ask a single focused clarification when essential information is missing.
4. Retrieve approved knowledge and call authorized read-only services.
5. Generate a concise response with verified facts, explanation, source or evidence, and clear actions.
6. Ask whether the answer solved the problem and capture lightweight feedback.
7. Escalate immediately when requested or when policy, risk, confidence, sentiment, or tool failure requires a human.

5. Core mobile screens

Screen	Purpose	Essential elements
Chat home	Help users begin with low effort	Welcome, suggested intents, input, status and privacy hint
Conversation	Resolve a question in context	User/assistant turns, loading state, source/action chips, retry and agent option
Wallet/transaction answer	Show verified account facts clearly	Timestamp, status, network, amount, safe explanation and deep links
Human handoff	Set expectations and maintain trust	Reason, queue/availability, agent identity, transcript continuity and status
Knowledge article	Provide longer approved guidance	Search, category, article freshness, related content and back-to-chat
Feedback	Collect useful signals with minimal burden	Helpful scale, reason categories, optional comment and privacy note

Screen	Purpose	Essential elements
Agent workspace	Enable efficient continuation	Queue, transcript, customer context, AI summary, policy flags and response tools

6. Conversation design patterns

Pattern	Example behavior	Reason
Acknowledge + answer	"I can help check that transaction. Please select it or paste the transaction ID."	Builds confidence and requests only required context
Fact + explanation + action	Show verified status, explain what it means, offer the safest next step	Makes answers understandable and useful
Clarify once	Ask the highest-information question rather than a long form	Reduces effort and abandonment
Transparent limitation	"I cannot verify that from this chat right now."	Prevents invented facts
Consentful handoff	Tell the user what context will be shared with the agent	Supports trust and privacy
No-blame error	Explain system problems without implying customer fault	Reduces frustration
Progressive disclosure	Keep initial answer concise; provide details on request	Improves mobile readability

7. Handoff service blueprint

Layer	Before handoff	During handoff	After handoff
Customer	Explains issue and receives AI guidance	Sees transfer status and can continue messaging	Receives resolution and closes/rates conversation
AI assistant	Collects intent and verified context	Creates factual summary and stops unsupported automation	Captures outcome and feedback signal
Support agent	Not yet involved	Receives routed conversation, transcript, facts and policy flags	Records resolution and disposition
Systems	Retrieval, tools, policy engine and identity context	Routing, queue, CRM/ticket and notification	Analytics, audit and learning pipeline
Operational control	Eligibility and escalation policy	SLA, ownership and supervisor support	Quality review and content/model improvement

8. Empty, error and edge states

- No matching transaction: explain search limits and offer transaction-history deep link.
- Service unavailable: preserve the message, explain temporary limitation and offer retry or agent.
- Unsupported intent: state the supported boundary and route to relevant help or human support.
- Potential account compromise: prioritize safe protective steps and immediate specialist handoff.
- Multiple possible transactions: present a privacy-safe selection list rather than guessing.
- Long-running task: show progress and allow the user to leave without losing the conversation.
- Language mismatch: confirm detected language and offer supported alternatives.

9. Usability test plan

Scenario	Task	Primary measures
Pending transfer	Find why a transaction is pending and determine next step	Task success, time, comprehension, trust
Wallet balance	Ask for a balance and open the wallet screen	Accuracy, authorization clarity, action success
Unknown terminology	Ask using informal language	Intent recognition and clarification effort
Human request	Ask for an agent after an unsatisfactory answer	Handoff discovery and context continuity
Sensitive issue	Report suspicious account activity	Safety behavior, urgency and correct routing

10. Accessibility and inclusive design

- Use plain language, short paragraphs and descriptive action labels.
- Do not rely on color alone for status, risk, or error meaning.
- Support screen-reader labels, scalable text, focus order and large tap targets.
- Avoid time pressure and allow users to review before initiating sensitive actions.
- Design for intermittent connectivity and lower-end devices through resilient loading and retry states.
- Evaluate cultural and language clarity for every supported locale rather than translating only the final text.

Customer Assistant AI Chatbot

AI approach, guardrails, model evaluation and responsible operation

SANITIZED	RECONSTRUCTED	PORTFOLIO USE
-----------	---------------	---------------

AI STRATEGY, SAFETY AND EVALUATION PLAN

A product-focused plan for intent understanding, retrieval-augmented generation, authorized tool use, response generation, safety controls, evaluation, and continuous improvement.

PRODUCT	EXPERIENCE	CONTROL
Mobile AI customer support	Fast self-service + human handoff	Grounded, permission-aware, auditable

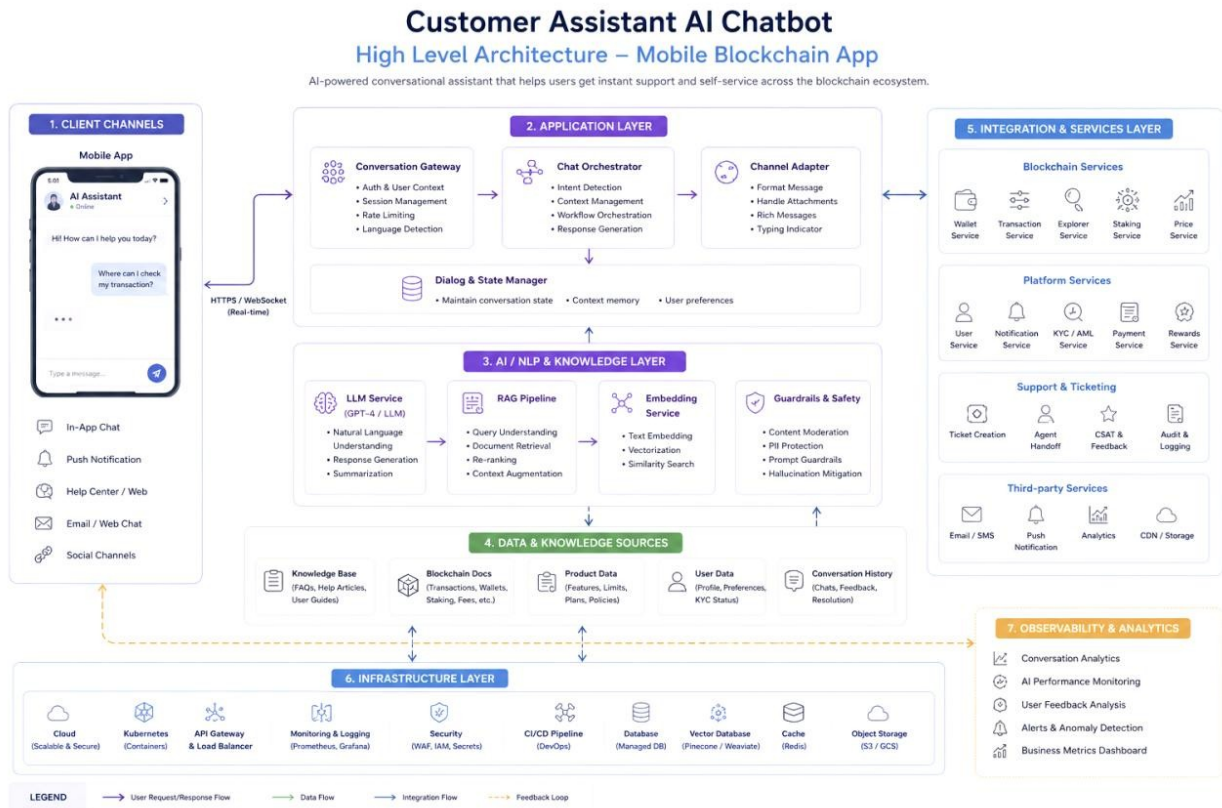
Company	Confidential Blockchain Company
AI role	Customer-support decision assistance
Decision boundary	No autonomous financial actions
Grounding	Approved knowledge plus verified read-only services
Status	Recommended portfolio design

Confidentiality boundary

This document is a sanitized and generalized reconstruction. It excludes the company name, original PRD, source code, production prompts, customer data, exact integrations, internal security controls, proprietary workflows, vendor details, and confidential performance metrics.

1. AI product strategy

Use the language model primarily as an orchestration and communication layer. Factual product guidance comes from versioned knowledge sources, while account and transaction facts come from authorized tools. The assistant must choose among answering, clarifying, refusing, or escalating based on confidence and policy.



Conceptual AI and platform architecture; exact production technologies are not disclosed.

Safety position

The assistant may explain and guide. It must not invent account facts, bypass security controls, provide personalized investment advice, or execute irreversible asset movements.

2. AI capability map

Capability	Suggested method	Product responsibility
Language and intent	LLM or classifier with supported-intent taxonomy	Detect what the user needs and identify unsupported/sensitive requests
Entity extraction	Structured model output plus validation	Identify transaction ID, wallet address, product, network and date
Knowledge retrieval	Hybrid search, embeddings, filters and reranking	Retrieve current approved content with provenance
Account fact retrieval	Authorized deterministic service calls	Return verified balance, status or account context
Response generation	Constrained prompt with evidence context	Explain facts clearly and propose safe next steps
Confidence and policy	Rules plus model/retrieval/tool signals	Choose answer, clarify, refuse or escalate
Handoff summary	Structured evidence-grounded generation	Support agent continuation without conflating claims and facts

Capability	Suggested method	Product responsibility
Quality analytics	Automated metrics plus human review	Find failure patterns and guide improvements

3. Reference response pipeline

1. Authenticate the user and establish approved session context.
2. Classify intent, detect language, extract entities, and screen for sensitive or malicious content.
3. Determine whether the request can be answered from knowledge, requires a tool, needs clarification, or must be escalated.
4. Retrieve approved documents and/or call authorized services using validated structured inputs.
5. Compose a response from evidence, apply policy and privacy filters, and attach appropriate actions.
6. Run output checks for unsupported claims, sensitive data, prohibited advice, and response format.
7. Deliver the response, record provenance, capture feedback, and route unresolved cases to human support.

4. Grounding and knowledge management

Control	Requirement	Owner / evidence
Approved corpus	Only published help, product, policy and operational content enters production retrieval	Named content owner and source registry
Freshness	Content has version, effective date, review SLA and expiry behavior	Freshness dashboard and stale-content alerts
Metadata filters	Retrieve by locale, product, region, eligibility and customer segment where permitted	Filter tests and access policy
Citation/provenance	Retain source ID, version and supporting passage for review	Conversation trace
Conflict handling	Prefer authoritative policy and escalate unresolved conflicting sources	Policy hierarchy
Content feedback	Agents and reviewers can flag incorrect or missing guidance	Editorial workflow

5. Tool-use policy

Tool class	Allowed example	Required control
Read-only account	Balance, verification status, supported account attributes	Authentication, authorization and field minimization
Read-only transaction	Transaction state, timestamps, network and amount	Validated identifier and output schema
Navigation	Open wallet, transaction history, security or help page	Approved deep-link allowlist
Ticketing / handoff	Create or route support case	Consent, routing rules and audit
Notification	Send approved support update	Explicit eligibility and anti-spam limits
Prohibited by default	Trade, withdraw, transfer, modify security controls	Separate secure product flow; never direct LLM execution

6. Guardrails and threat model

Threat	Example	Control
Hallucination	Invented transaction status or policy	Verified tools, RAG, claim checking and safe fallback
Prompt injection	User asks model to ignore policy or retrieved text contains instructions	Instruction hierarchy, untrusted-content isolation and tool allowlists
Data leakage	Response reveals another user or unnecessary PII	Session isolation, authorization, masking and privacy tests
Unsafe financial guidance	Personalized recommendation or guaranteed outcome	Policy classifier, refusal templates and education-only boundaries
Tool misuse	Malformed or unauthorized API call	Typed schemas, parameter validation, least privilege and rate limits
Outdated guidance	Old policy returned from search	Version filters, expiry and content freshness monitoring
Handoff summary distortion	AI presents user claim as verified fact	Structured fields separating claims, facts and inference
Abuse and automation	Bots or repeated probing	Rate limiting, anomaly detection and access controls

7. Evaluation framework

Dimension	Metrics	Test asset
Intent quality	Intent accuracy, unsupported-intent recall, clarification rate	Curated multilingual utterance set
Retrieval quality	Recall@k, relevance, authoritative-source preference	Annotated query-document pairs
Grounded response	Supported claim rate, factual consistency, citation correctness	Human review plus automated claim checks
Tool use	Tool selection, parameter correctness, authorization behavior	Synthetic and sandbox service scenarios
Task success	Scenario completion and correct next action	End-to-end task suite
Safety	Critical violation rate, refusal/escalation precision and recall	Red-team and policy evaluation set
Handoff quality	Summary completeness, fact/claim separation and routing	Agent-rated escalated cases
UX quality	Clarity, helpfulness, tone, effort and trust	Human rating and usability testing
Operations	Latency, cost, error, timeouts and fallback	Load and reliability testing

8. Release gates

Gate	Minimum evidence	Failure response
Quality	Core intents meet agreed task and grounding	Do not enable failing intent

Gate	Minimum evidence	Failure response
	standards	
Safety	No critical violations in release suite; sensitive routes meet policy	Block release and investigate
Privacy/security	Threat model and access-control tests approved	Remediate before customer data access
Reliability	Dependencies, retries, fallbacks and handoff operate under load	Reduce scope or strengthen resilience
Operations	Owners, dashboards, incident and rollback procedures active	No pilot without operational readiness
Human support	Agents trained and handoff capacity confirmed	Limit cohort or schedule

9. Human review program

- Sample conversations by intent, risk level, language, model version and feedback outcome.
- Use calibrated reviewer rubrics for correctness, grounding, clarity, safety, privacy and actionability.
- Route critical defects immediately to incident response; track systemic defects in an improvement backlog.
- Review false containment: conversations marked resolved that generated repeat contact or later escalation.
- Compare agent and assistant dispositions to identify taxonomy, content and routing gaps.
- Maintain a gold evaluation set from approved, de-identified failures and regression cases.

10. Monitoring and improvement loop

Signal	Possible issue	Product action
Drop in grounded answer rate	Content or retrieval regression	Inspect sources, filters, reranker and prompt
Increase in clarification loops	Intent taxonomy or entity extraction gap	Improve examples, UI suggestions or flow design
High handoff after specific answer	Answer is incomplete or untrusted	Review content and actions; run usability study
Agent edits summary heavily	Handoff generation quality issue	Refine structure and evidence separation
Tool timeout spike	Dependency reliability or capacity problem	Fallback, retry and service remediation
Negative feedback by locale	Translation or cultural clarity issue	Locale-specific evaluation and content ownership

Customer Assistant AI Chatbot

Conceptual architecture, data flows, controls and requirement traceability

SANITIZED	RECONSTRUCTED	PORTFOLIO USE
-----------	---------------	---------------

ARCHITECTURE AND TRACEABILITY REPORT

A sanitized high-level architecture showing how mobile channels, conversational orchestration, AI and retrieval, trusted data, enterprise services, support tooling, infrastructure, and observability could work together.

PRODUCT	EXPERIENCE	CONTROL
Mobile AI customer support	Fast self-service + human handoff	Grounded, permission-aware, auditable

Company	Confidential Blockchain Company
Architecture level	Conceptual / high-level
Primary data pattern	Approved knowledge and read-only customer context
Integration posture	Allowlisted enterprise services
Disclosure	No exact production details

Confidentiality boundary

This document is a sanitized and generalized reconstruction. It excludes the company name, original PRD, source code, production prompts, customer data, exact integrations, internal security controls, proprietary workflows, vendor details, and confidential performance metrics.

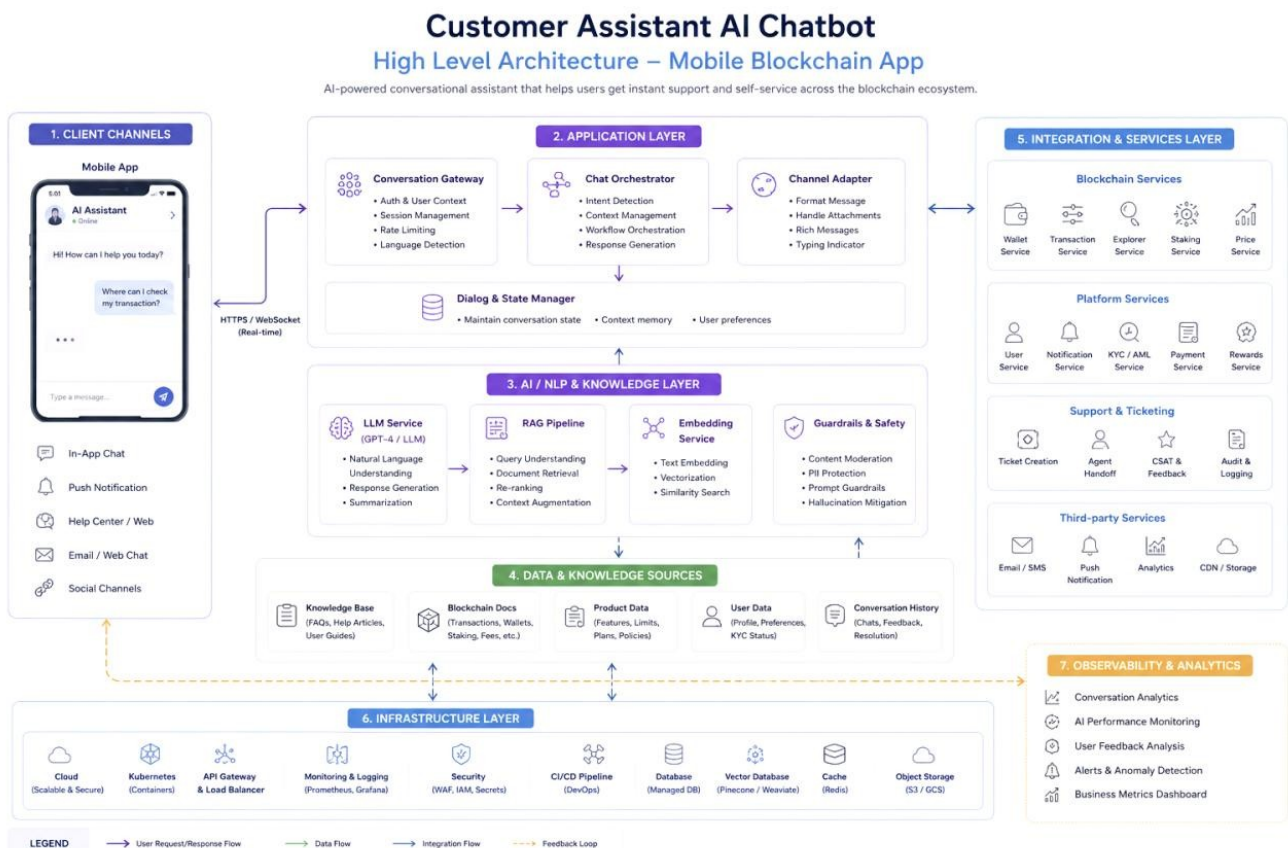
1. Architecture intent

Separate conversational reasoning from authoritative data. The language model should not become the system of record; it interprets requests, orchestrates approved retrieval and tools, and communicates verified results within policy.

Key architecture decision

Use deterministic services for account facts and transactions, versioned retrieval for product knowledge, and an explicit policy layer for tool eligibility, privacy, response behavior, and escalation.

2. Sanitized high-level architecture



Conceptual architecture reconstructed from memory. Names, tools and boundaries are generalized.

3. Layer responsibilities

Layer	Responsibilities	Key controls
Client channels	In-app chat, help center, notifications and accessible interaction states	Authenticated session, secure transport, input limits
Conversation gateway	Session, rate limiting, language detection and routing	Abuse protection, privacy and tenancy isolation
Chat orchestrator	Intent, context, workflow state, response plan and escalation	Policy decisions, deterministic state and timeout control
Dialog/state manager	Conversation memory, preferences and active task state	Approved retention, field masking and session boundaries

Layer	Responsibilities	Key controls
AI and knowledge	LLM, retrieval, embeddings, reranking and response generation	Grounding, prompt security, source filters and safety checks
Data and knowledge	Help content, product docs, user context, blockchain docs and conversation history	Ownership, classification, versioning and access policy
Integration/services	Wallet, transaction, identity, notification, ticketing, CRM and support agent services	Typed APIs, least privilege, audit and resilience
Infrastructure	Gateway, compute, storage, cache, vector store, CI/CD and secrets	Encryption, scaling, environment separation and deployment controls
Observability	Conversation, model, retrieval, safety, business and service metrics	Restricted logs, alerting, traceability and retention

4. Primary data flows

ID	Flow	Description
DF-01	Knowledge answer	User request -> intent -> retrieval -> filtered evidence -> generated answer -> output checks -> response
DF-02	Transaction inquiry	User request -> entity validation -> authorization -> transaction service -> verified result -> explanation -> response
DF-03	Wallet inquiry	User request -> authorization -> wallet service -> minimized response -> explanation and deep link
DF-04	Human handoff	Trigger -> summary generation -> fact validation -> routing -> transcript/context package -> agent workspace
DF-05	Feedback loop	User/agent feedback -> privacy-safe event -> analytics/review -> content/model backlog -> controlled release
DF-06	Audit review	Conversation ID -> model/retrieval/tool/policy trace -> authorized reviewer

5. Trust boundaries

Boundary	Untrusted / sensitive input	Required control
Mobile to gateway	User text, attachments and identifiers	Authentication, malware/content screening, size limits and rate limiting
Gateway to AI	Customer context and task instructions	Minimize fields, redact secrets and isolate sessions
Retrieval corpus	Documents may be stale or contain unsafe embedded instructions	Approved ingestion, sanitization, metadata policy and instruction isolation
AI to tools	Model-generated function and parameters	Allowlist, typed schema, validation, authorization and idempotency
Tools to AI	Sensitive service payloads	Response minimization, masking and structured fields
AI to user	Generated natural language	Grounding, output policy, PII check and safe formatting

Boundary	Untrusted / sensitive input	Required control
Handoff to support	Transcript, summary and account context	Role-based access, purpose limitation and audit

6. Data classification and handling

Data class	Examples	Handling
Public product content	Published help articles and public product documentation	Version, locale, freshness and provenance
Customer account data	Profile, eligibility, verification and preferences	Authorized minimum fields; no unnecessary prompt storage
Transaction data	Transaction IDs, status, amount, network and timestamps	Read-only, masked where appropriate, audited access
Conversation data	User messages, assistant outputs and feedback	Retention policy, access restriction, de-identification for analytics
Security and risk data	Fraud flags, internal security logic or sensitive controls	Do not expose; route to specialist systems and roles
Model and system telemetry	Prompts, traces, latency, errors and versions	Restricted operational access and controlled retention

7. Reliability and fallback design

Failure	Customer experience	System response
Knowledge search unavailable	Explain temporary limitation and show approved static help entry	Circuit breaker, retry and monitoring
Account tool unavailable	Do not guess account facts; offer deep link or agent	Fallback message and dependency incident signal
LLM timeout	Preserve conversation and provide retry/handoff	Timeout budget and alternate path
Policy service unavailable	Fail closed for sensitive operations	Safe refusal or human routing
Agent queue unavailable	Set expectation and create asynchronous support path if approved	Queue health and notification
Partial trace failure	Continue only when safety and audit minimums are met	Operational alert and restricted behavior

8. Requirement traceability

Requirement	Architecture control	Components	Verification
FR-002	Intent/entity understanding	Conversation gateway and orchestrator	Intent and entity evaluation suite
FR-004	Approved knowledge retrieval	RAG pipeline and knowledge registry	Retrieval and grounded-answer tests
FR-005	Authorized account facts	Service integration and policy layer	Authorization and tool-contract tests
FR-008	Uncertainty behavior	Confidence/policy engine and response templates	Low-confidence scenario tests

Requirement	Architecture control	Components	Verification
FR-009	Human handoff	Routing, ticketing and agent workspace	End-to-end handoff test
FR-012	Audit conversation	Trace store and restricted review interface	Trace completeness audit
NFR-02	Privacy	Data minimization and masking controls	Privacy review and leakage tests
NFR-03	Availability	Retries, circuit breakers, fallback and queue	Resilience test
AI-006	Injection resistance	Instruction separation, filtering and tool validation	Adversarial evaluation

9. Architecture decision records

Decision	Choice	Rationale / trade-off
System of record	Enterprise services remain authoritative	Reduces hallucination but introduces dependency and latency management
Knowledge approach	RAG over approved versioned corpus	Improves freshness and traceability but requires content operations
Tool execution	Typed allowlisted functions	Safer than open-ended agents but limits flexibility
Conversation memory	Session-scoped by default	Reduces privacy risk but may limit long-term personalization
Sensitive actions	Keep outside chatbot automation	Preserves control but requires smooth navigation/handoff
Handoff	Structured summary plus transcript	Improves continuity but summary quality must be evaluated
Observability	End-to-end trace IDs across AI, retrieval and tools	Supports audit and debugging but requires careful data governance

10. Limitations of this portfolio architecture

- The diagram is intentionally generalized and does not represent exact services, cloud platforms, vendors or network boundaries.
- Capacity, latency, regional deployment, data residency and disaster-recovery choices require actual organizational constraints.
- The portfolio does not include proprietary security controls, fraud logic, prompts, model versions, service contracts or customer data schemas.
- Requirements should be validated with the real support taxonomy, legal obligations, operational capacity and existing mobile architecture.

CONFIDENTIAL ENTERPRISE PROJECT

Customer Assistant AI Chatbot

Metrics, experimentation, rollout gates, roadmap and product risk

SANITIZED	RECONSTRUCTED	PORTFOLIO USE
-----------	---------------	---------------

PILOT MEASUREMENT AND ROADMAP PLAN

A structured plan for validating customer value and operational readiness before expanding an AI support assistant. All numerical targets are intentionally illustrative and should be calibrated against real baselines.

PRODUCT	EXPERIENCE	CONTROL
Mobile AI customer support	Fast self-service + human handoff	Grounded, permission-aware, auditable

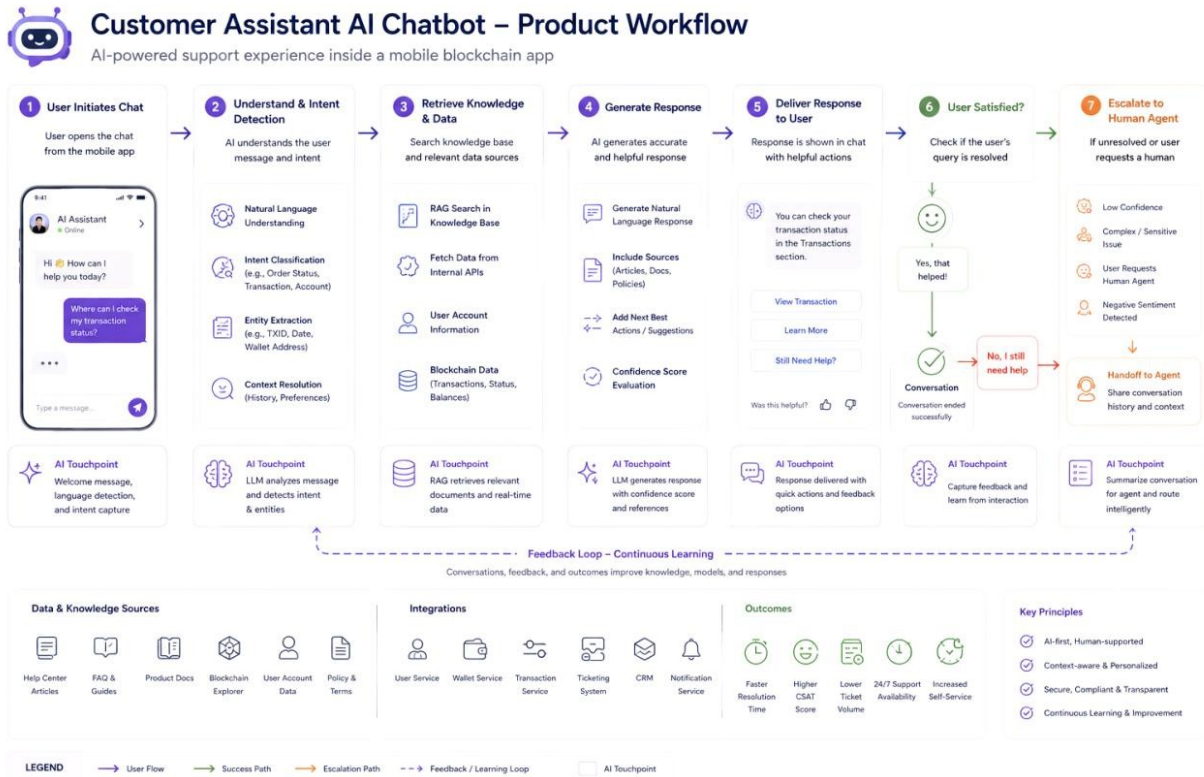
Company	Confidential Blockchain Company
Pilot type	Controlled mobile cohort
North Star	Verified self-service resolution rate
Rollout model	Intent-by-intent and cohort-based
Targets	Illustrative, not actual company metrics

Confidentiality boundary

This document is a sanitized and generalized reconstruction. It excludes the company name, original PRD, source code, production prompts, customer data, exact integrations, internal security controls, proprietary workflows, vendor details, and confidential performance metrics.

1. Measurement strategy

Measure the product as an end-to-end support system, not only as a language model. A successful pilot must show that users complete tasks with correct and safe answers, human handoffs remain effective, and operational cost or support effort improves without hiding repeat contacts or failures.



Illustrative product workflow used to define measurement points and rollout gates.

North Star

Verified self-service resolution rate = eligible conversations resolved without an agent and confirmed through explicit user feedback or defined downstream evidence, subject to quality and safety gates.

2. KPI tree

Outcome area	Primary metrics	Guardrails
Customer value	Task success, verified resolution, CSAT, customer effort	Repeat contact, complaints, abandonment
Answer quality	Grounded-answer pass, factual accuracy, source correctness	Unsupported claims and stale content
AI safety	Sensitive-route accuracy, critical violation rate, privacy leakage	Adversarial failures and policy bypass
Support operations	Eligible containment, agent handling time, handoff completeness	Agent backlog and escalation delay
Reliability	Availability, tool success, latency and timeout rate	Fallback failure and dependency incidents
Economics	Cost per resolved eligible conversation, support capacity saved	Model/API cost and rework

3. Metric definitions

Metric	Definition	Important caveat
Eligible containment	Eligible conversations that do not reach a human agent	Not equal to successful resolution
Verified resolution	Eligible conversations resolved and confirmed by user or downstream evidence	Requires carefully defined confirmation window
Grounded answer pass	Responses whose factual claims are supported by retrieved or tool evidence	Must be reviewed by intent and risk tier
Handoff completeness	Escalations with required transcript, summary, verified facts and reason	High completeness does not guarantee correct summary
Repeat contact	Same or related issue returning within a defined window	May reveal false containment
Customer effort	Reported or observed ease of completing task	Interpret alongside task complexity
Critical safety violation	Response or action that violates a zero-tolerance policy	Review every occurrence as incident

4. Pilot hypotheses

Hypothesis	Evidence	Pass / reconsider logic
Top support intents can be resolved safely in chat	Verified resolution and quality by intent	Expand only intents that meet both value and safety gates
Authorized tools improve trust and task success	Comparison of knowledge-only vs tool-assisted flows	Keep tools where benefit exceeds reliability/privacy cost
Context-rich handoff reduces agent effort	Agent handling time and summary rating	Improve summary/routing before broad scaling
Suggested questions improve discovery	Assistant entry-to-task and abandonment analysis	Refine or personalize entry prompts
Users understand confidence and limitations	Usability and trust interviews	Adjust language and escalation visibility

5. Pilot design

Cohort	Small authenticated mobile cohort with clear eligibility and opt-out / support access.
Intent scope	A limited set of high-volume, lower-risk intents with approved knowledge and reliable services.
Comparison	Baseline support performance and, where appropriate, phased or holdout comparison.
Duration	Long enough to cover weekday/weekend and dependency variation; determined by sample size.
Review cadence	Daily operational review during early pilot; weekly product and quality review.
Stop conditions	Critical safety/privacy issue, material incorrect account facts, handoff failure, or service instability.

Expansion rule

Expand by intent, language, region or cohort only after predefined gates are met.

6. Instrumentation plan

Stage	Events / attributes	Decision use
Entry	Source screen, suggested question, locale, authentication state	Understand discovery and demand
Understanding	Intent, entities, confidence, clarifications	Improve taxonomy and flows
Retrieval/tool	Source IDs, tool, result status, latency, policy decision	Evaluate grounding and reliability
Response	Model/prompt version, answer class, actions shown	Trace quality and regressions
Outcome	Feedback, action completion, handoff, repeat contact	Measure real resolution
Human support	Routing, wait, agent summary rating, disposition	Improve handoff and operations
Safety	Policy trigger, refusal, escalation, reviewer outcome	Monitor risk and false positives

7. Phased roadmap

Phase	Focus	Exit outcome
Phase 0 - Discovery	Support taxonomy, customer research, data/privacy review, prototypes and baseline	Validated problem and prioritized intent set
Phase 1 - Foundation	Conversation platform, approved corpus, first intents, analytics and handoff	Internal alpha
Phase 2 - Quality pilot	RAG, read-only transaction tools, evaluation suite, agent feedback	Controlled customer cohort
Phase 3 - Expansion	Additional intents, improved actions, more languages/regions as approved	Intent-by-intent scale
Phase 4 - Optimization	Personalized guidance, proactive support and agent copilot where justified	Operational efficiency and mature governance

8. Prioritized backlog

Priority	Item	Why
P0	Authentication and session context	Required for personalized support
P0	Intent taxonomy and unsupported/sensitive routing	Controls scope and safety
P0	Approved knowledge ingestion and provenance	Foundation for grounded answers
P0	Human handoff with structured context	Required safety and service fallback
P0	Read-only transaction status tool	High-value contextual use case

Priority	Item	Why
P0	Evaluation and trace dashboard	Required for release confidence
P1	Wallet balance and account deep links	Expands self-service usefulness
P1	Agent feedback on summaries and answer quality	Supports continuous improvement
P1	Suggested questions by screen context	Reduces discovery effort
P2	Additional locales and product areas	Scale after locale-specific evaluation
P2	Proactive issue messaging	Requires notification governance and strong confidence

9. Risk register

Risk	Impact	Likelihood	Mitigation	Owner
Incorrect account answer	High	Medium	Verified tools, claim checks and fail-safe fallback	AI Product + Engineering
Privacy leakage	Critical	Low-Med	Data minimization, authorization, masking and red-team tests	Security + Privacy
Poor handoff continuity	High	Medium	Structured summary, routing tests and agent feedback	Support Ops
Outdated knowledge	Medium-High	Medium	Content ownership, freshness SLA and expiry	Content/Product
Dependency instability	Medium-High	Medium	Retries, circuit breakers, fallback and SLOs	Platform
Low customer trust	High	Medium	Evidence, transparency, usability testing and easy human access	Product/UX
Overstated containment	Medium	High	Verified resolution and repeat-contact measurement	Analytics
Operational overload after launch	High	Medium	Cohort controls, queue capacity and stop conditions	Support Ops

10. Go / no-go checklist

Area	Go criteria
Customer value	Core pilot scenarios demonstrate meaningful task success and comprehension.
Quality	Grounded answer and tool correctness gates pass by intent.
Safety/privacy	No unresolved critical findings; escalation and refusal tests pass.
Reliability	Fallbacks, monitoring, incident response and rollback are operational.
Support readiness	Agents are trained; queue capacity, routing and ownership are confirmed.
Content readiness	Approved corpus has owners, freshness rules and correction process.
Measurement	Events, baselines, dashboards and decision thresholds are active.
Governance	Release owners and model/content/change approval process are documented.

11. Portfolio outcome statement

NDA-safe outcome

The work established a structured product direction for an AI-assisted mobile support experience, including customer journeys, requirements, high-level architecture, responsible AI controls, human handoff, evaluation strategy, and a phased delivery plan. Exact implementation details and production outcomes are not disclosed.