

Mock Exam A

UXQCC Certified Professional in Software Usability with GenAI – Foundation Level

This mock exam is a practice exam for preparation. It is NOT the actual certification exam, but it follows its rules per the exam regulations: 40 multiple-choice questions with four answer options each, of which exactly one is correct. Each correct answer counts one point; there is no negative marking. Working time: 60 minutes (for an exam taken in a foreign language, 15 minutes of additional time may be granted). Passed from 26 of 40 points (65 percent). The composition follows the chapter and K-level matrix of the exam regulations: 30 questions at levels K1/K2 and 10 questions at level K3; the K level is stated for each question. Recommendation: Work through Part 1 under exam conditions without aids and only afterwards evaluate with Part 2.

Part 1: Exam Questions

Question 1 (K1): In a development team, the terms usability, user experience, and UX quality are being discussed. Which mapping renders the terms correctly?

- A) Usability refers to the speed of a system, user experience to freedom from errors, and UX quality to compliance with the style guide.
- B) Usability concerns exclusively operation by mouse and keyboard, while user experience is relevant only for mobile applications; UX quality combines both for marketing purposes.
- C) Usability refers to the extent to which specific users can achieve their goals effectively, efficiently, and with satisfaction in a specific context of use; user experience additionally covers the entire usage experience before, during, and after use; UX quality refers to the quality of use and usage experience of a system.
- D) Usability is the older technical term, which has been fully superseded by user experience; UX quality refers to the visual modernity of a product.

Question 2 (K2): A Scrum team works with discipline: The backlog is well maintained, the Definition of Done is observed, every increment is shown in the sprint review. After several releases, users nevertheless complain about incomprehensible flows. Which explanation fits best?

- A) The team presumably did not carry out the Scrum events in the intended order, which caused usability aspects to get lost.
- B) Sprint reviews are the wrong format for user topics; the complaints could have been avoided through additional retrospectives.
- C) The Definition of Done was too strict and prevented the team from responding flexibly to user needs.
- D) Agile processes structure the work, but they do not automatically ensure that user goals, context of use, system states, error cases, and understandability are substantively worked on and verified.

Question 3 (K2): A team observes: For a standard login form, GenAI delivers very usable suggestions right away. For a domain application controlling laboratory processes, however, the suggestions appear generic and partly unsuitable. Which explanation corresponds to the basic rule of the syllabus?

- A) The AI was presumably trained on too few form examples; with better prompts, both results would be equally good.
- B) Login forms are technically simpler than laboratory processes, which is why the suggestions are better there.
- C) The AI favors consumer applications, because domain applications are economically less relevant.
- D) The login form follows a widespread, well-documented pattern that the AI knows from many examples; the laboratory context with its domain logic, roles, and risks is unfamiliar context for the AI, which the team must supply and verify.

Question 4 (K1): In a workshop, participants are to distinguish central basic usability terms from other technical terms. Which series of terms belongs entirely to the basic usability vocabulary from Chapter 2?

- A) Persona, stakeholder map, customer journey, value proposition, north star metric, and product vision.
- B) Responsiveness, latency, cache, rendering, frame rate, and compression.
- C) Retrospective, timebox, increment, product goal, commitment, and Definition of Ready.
- D) User goal, context of use, mental model, conformity with user expectations, cognitive load, visual hierarchy, and accessibility.

Question 5 (K2): In a warehouse management system, pickers regularly overlook a warning highlighted in red, although it is large and centrally placed. They work under time pressure, scan items every few seconds, and mostly look at the input field while doing so. Which explanation describes the phenomenon best?

- A) Human perception is selective, limited, and context-dependent: Under time pressure and with a focused task, even conspicuous things are overlooked when they lie outside the focus of attention or do not match expectations.
- B) The color red is unsuitable for warnings, because it has a calming effect in the work context and is therefore ignored.
- C) The pickers perceive the warning completely, but consciously decide not to heed it; that is a discipline problem and not a perception issue.
- D) The warning is overlooked because screens in warehouses are generally too small; on larger displays the problem would not occur.

Question 6 (K2): A booking system uses the button “Close case”. Some users expect that this completes and saves the case; others fear that their inputs will be discarded and therefore never click it. Why are mental models, conformity with user expectations, and cognitive load the right explanatory framework for this problem?

- A) Because the three concepts describe how an interface needs as few clicks as possible and thereby becomes more efficient.
- B) Because users interpret the label against the background of their idea of the system: If the system behavior does not match the expectation, uncertainty, operating errors, and additional mental effort arise with every decision.
- C) Because the three concepts mainly determine what visual design a button must have to be recognized as clickable.
- D) Because mental models play a role only for first-time users, and the problem therefore disappears by itself with growing experience.

Question 7 (K2): A GenAI writes about a dashboard draft: “The filter bar will be understood intuitively by most users.” How is this statement to be classified?

- A) As a hypothesis: The statement sounds like a factual claim about user behavior, but it does not rest on observation of real users and must be checked professionally or empirically before it carries decisions.
- B) As a reliable finding, because the AI was trained on patterns from very many comparable dashboards.
- C) As user feedback, because the AI takes the users' perspective and predicts their understanding.
- D) As an irrelevant statement that needs neither checking nor documentation, because it is worded positively.

Question 8 (K3): A team has an AI assess the registration flow. The AI answers at length, names correct technical terms, and concludes: “The flow complies with common usability standards in all respects.” A developer points out that the flow contains a mandatory field whose purpose no one in the team itself can explain. Which reaction shows the right critical handling of the AI answer?

- A) Follow the blanket approval, because the AI uses technical terms correctly and thereby demonstrates professional competence.
- B) Ask the AI again and, if the answer is the same, adopt the assessment, because repetition increases reliability.
- C) Discard the answer and refrain from AI assessments in the future, because the AI did not object to the mandatory field.
- D) Question the blanket statement: A comprehensive quality judgment without knowledge of context of use and domain logic is seemingly plausible; check the unclear mandatory field in a targeted way and ask the AI for concrete, justified individual assessments instead of blanket judgments.

Question 9 (K1): A product owner creates a term overview for the team. Which series contains exclusively the central discovery terms from the learning objective?

- A) User, user goal, task, context of use, problem hypothesis, and product assumption.
- B) Epic, feature, task, subtask, spike, and bug.
- C) Click rate, bounce rate, dwell time, session length, conversion rate, and retention.
- D) User, customer, buyer, licensee, client, and stakeholder.

Question 10 (K2): A stakeholder wishes for “a chatbot like the competitor's” and wants to start implementation immediately. The team proposes first clarifying users, tasks, and context of use. Why is this intermediate step professionally required?

- A) Because a chatbot is technically complex and the team must first clarify the architecture questions before user topics become relevant.
- B) Because stakeholder wishes should generally be rejected until a complete market analysis is available.
- C) Because otherwise the solution idea stands before the problem: Only the understanding of users, tasks, and context shows which problem is actually to be solved and whether a chatbot would even be a suitable means for it.
- D) Because without documented context of use, no user stories may be added to the backlog.

Question 11 (K2): A team has scattered information: interview notes from two customer conversations, an Excel list with support hints, app store reviews, and some observations from a trade fair demo. Which description captures the sensible role of GenAI in processing this material?

- A) GenAI weighs the sources against each other and decides which of them are credible enough to be considered in the further process.
- B) GenAI brings the heterogeneous sources together, summarizes them, makes recurring topics visible across the sources, and formulates initial hypotheses from them, which the team then checks.
- C) GenAI supplements the incomplete interview notes analogously, so that a complete picture of user needs emerges.
- D) GenAI assesses which of the named user groups is most valuable for the product and aligns the summary with that.

Question 12 (K3): After analyzing analytics data, an AI formulates: “Users abandon the registration process because the form is too long.” How does the team formulate a correct hypothesis from this and distinguish it from reliable findings?

- A) “We suspect that the form length contributes to abandonments in registration. So far, only the abandonment itself is proven; we will check the cause through observation, user inquiry, or a comparison test.”
- B) “The form length causes the abandonments. This is proven by the analytics data and can be documented as a reliable finding.”
- C) “The AI has identified the cause; we document it as a finding as long as no contradicting data appears.”
- D) “Since causes cannot be determined reliably anyway, we dispense with a hypothesis and shorten the form as a precaution.”

Question 13 (K1): In a team training, the basic terms of usability-related backlog refinement are to be named. Which series contains exactly the terms named in the learning objective?

- A) Persona, scenario, journey map, empathy map, and prototype.
- B) Usability requirement, user story, acceptance criterion, Definition of Ready, and Definition of Done.
- C) Test case, test plan, test coverage, regression test, and acceptance test.
- D) Product goal, sprint goal, increment, commitment, and timebox.

Question 14 (K2): In a refinement, a developer says: “Requirements are the business department's job; we implement what is specified. Usability can always be improved afterwards.” Which reply best justifies why usability requirements belong in requirements work?

- A) Usability requirements belong in requirements work because they relieve the business department and make the specification shorter.
- B) Usability requirements belong in requirements work because they prevent the business department from expressing its own design wishes.
- C) Usability requirements belong in requirements work because subsequent improvements are legally impermissible.
- D) Whether users can cope with a function, interpret it correctly, and remain able to act even in the error case is a quality requirement like security or performance; if it is only considered afterwards, structure, flows, and states are often already so fixed that corrections become expensive.

Question 15 (K2): A team revises its working agreements. Which allocation of usability-relevant criteria to Definition of Ready and Definition of Done is professionally correct?

- A) DoR: The visual fine-tuning has been approved. DoD: The user goal of the story is known.
- B) DoR: All usability tests are completed. DoD: The open assumptions are documented.
- C) DoR: User goal and context of use are clarified, relevant system states and error cases named, open assumptions documented. DoD: Usability acceptance criteria fulfilled, error messages understandable, accessibility basics checked.
- D) DoR and DoD should contain no usability criteria, because these lie in the responsibility of a separate quality assurance.

Question 16 (K3): For the story “As an employee, I want to submit a travel expense report so that my expenses are reimbursed”, acceptance criteria with a usability focus are to be added. Which criterion achieves this best?

- A) “If a mandatory receipt is missing or an amount is implausible, the system shows understandably before submission which field is affected and how it can be corrected; after submission, the status of the report is visible at any time; all input fields are reachable by keyboard.”
- B) “The report is archived as PDF/A and transferred to the accounting system via the interface.”
- C) “The submission is processed in at most 15 seconds of server time and creates an entry in the audit log.”
- D) “The form complies with the corporate design and uses the approved typeface.”

Question 17 (K3): For a password reset story, an AI delivers eight acceptance criteria. Among them: “The link is valid for 24 hours”, “The page loads in under one second”, and “If the email address does not exist, the user receives a notice that no account exists”. How does the team handle the suggestions properly?

- A) Adopt all eight criteria, because they are worded concretely and verifiably and the AI evidently possesses security knowledge.
- B) Assess each criterion individually: adopt or adapt the sensible ones (for example, the validity period per the company's own security policy), and exclude the criterion with the notice about non-existent accounts with justification, because it discloses the existence of accounts and thereby creates a security and privacy problem.
- C) Reject the criteria wholesale, because with security-related functions no AI suggestions may be used, and write completely own criteria.
- D) Adopt the criteria unchanged and forward the security questions to the security team, which conducts a review after the release.

Question 18 (K1): A Scrum Master collects a list of typical usability tasks in implementation for the sprint board. Which list corresponds to the tasks named in the learning objective?

- A) Persona workshop, market study, competitor analysis, brand workshop, and pricing.
- B) Writing unit tests, merging code, fixing the pipeline, backing up the database, and rotating logs.
- C) Usability check, design critique, system state check, form test, accessibility basic check, and preparation of review questions for review formats.
- D) Team-building event, vacation cover, onboarding of new colleagues, and budget planning for the next quarter.

Question 19 (K2): In a team without a dedicated UX role, a tester says: “Checking understandability is not my job, I test against the specification.” The product owner says: “We have no one for user questions.” Which clarification describes the role contributions to UX quality correctly?

- A) The tester is right: Understandability is exclusively the job of UX specialists; without this role, the check is dropped without replacement.
- B) The product owner should transfer the responsibility for UX quality to the Scrum Master, because they are responsible for the team processes.
- C) UX quality should be awarded as its own work package to an external agency so that the team can concentrate on functionality.
- D) UX quality is a shared team task: The product owner clarifies user goals and priorities, the developers deliberately implement states and system feedback, testers also check behavior and understandability, the Scrum Master promotes the visibility of the tasks, and UX specialists add deeper expertise where needed.

Question 20 (K2): After a draft discussion, a developer says with disappointment: “Everyone just said what they liked or disliked. In the end, the loudest colleague won.” How does a design critique differ from this discussion?

- A) A design critique assesses drafts in a structured way against criteria agreed in advance, such as user goal, understandability, states, consistency, and accessibility; contributions refer to criteria instead of personal taste, and results are recorded.
- B) A design critique differs above all in that only design-experienced people may speak and the remaining team members listen.
- C) A design critique avoids conflicts by submitting criticism exclusively in writing and anonymously as a matter of principle.
- D) A design critique always ends with a binding approval by the majority of those present.

Question 21 (K1): A developer prepares a short presentation on the basic terms of UI design. Which series of terms covers exactly the basic terms named in the learning objective?

- A) Color harmony, golden ratio, moodboard, illustration style, visual language, and logo variants.
- B) Visual hierarchy, layout, spacing, typography, consistency, interaction, state, and control.
- C) Responsiveness, loading time, frame rate, caching, compression, and image optimization.
- D) Contrast ratio, alt text, screen reader, landmark, tab order, and ARIA attribute.

Question 22 (K1): A team plans prototypes for different questions. Which mapping of prototype type and purpose is correct?

- A) Low-fidelity prototype: final acceptance by stakeholders; high-fidelity prototype: first rough idea collection; horizontal prototype: in-depth examination of a single function; vertical prototype: overview of all areas.
- B) Scenario prototype: checking the server load; horizontal prototype: testing the data migration; wireframe: checking the brand effect; mockup: measuring access numbers.
- C) Low-fidelity prototype: early structure and variant discussion with little effort; high-fidelity prototype: realistic representation including visual design; horizontal prototype: showing the breadth of areas and navigation; vertical prototype: working out one selected flow in depth; scenario prototype: making a concrete usage case playable.
- D) Wireframe, mockup, and clickable prototype are three names for the same artifact in different project phases.

Question 23 (K2): A team observes: For the standard search bar with filters, the AI delivers solid drafts right away; for the domain-specific procurement review workflow, however, hardly usable ones. A colleague asks why that is. Which explanation is correct?

- A) The AI was evidently trained predominantly on search masks and is generally unsuitable for other UI types.
- B) Search bars are the only UI pattern the AI masters; for all other patterns, manual drafts are necessary.
- C) The difference is random; with renewed generation, the results would presumably be reversed.
- D) Search and filters belong to the frequent, well-documented UI patterns from which AI can derive typical structures and variants; the procurement review workflow, by contrast, is domain-specific and requires context and domain knowledge that the AI lacks.

Question 24 (K3): An AI delivers a draft for a terminal application at self-service kiosks in a citizens' office: small type, light gray labels, symbols with double meanings, and a confirmation dialog that reacts differently depending on the procedure. Which assessment applies the criteria from LO 6.6 correctly?

- A) Measured against user goal and context of use, the draft has serious weaknesses: low perceivability due to weak contrasts and small type, high cognitive load due to ambiguous symbols, violations of expectable behavior due to the inconsistent dialog, and accessibility problems for the broad, heterogeneous user group of a citizens' office.
- B) The draft is suitable, because light gray tones and small type look modern and the application thereby appears contemporary.
- C) The draft should be adopted, provided it is kept in the administration's corporate design and technically covers all procedures.
- D) The draft cannot be assessed before click numbers from productive operation are available.

Question 25 (K3): An AI generates an appealing draft for a document approval in a customer portal. While going through it, the team checks specifically for typical weaknesses of AI-generated drafts. Which checklist covers the points named in LO 6.8?

- A) Loading time of the preview images, size of the installation file, memory consumption, and compatibility with older browsers.
- B) Are loading, empty, and error states provided for? What happens with a missing approval permission, expired documents, or abandonment during the procedure? Are alternative usage situations such as deputy arrangements considered, and is accessibility taken into account?
- C) Does the color world correspond to the brand identity, are the icons drawn uniformly, and does the draft appear higher-quality than the competitor's?
- D) Can the draft be imported into Figma, are the layers cleanly named, and is the grid built consistently?

Question 26 (K1): A development team assembles its tool landscape for UI-related development. Which list describes AI-assisted coding tools with their typical fields of application correctly?

- A) Selenium and Cypress generate complete UI components via AI; Docker and Kubernetes check their accessibility.
- B) Figma and Sketch are AI coding tools, because their drafts are automatically translated into production-ready code.
- C) Tools like GitHub Copilot or Cursor support writing and completing UI components and

frontend code, generating tests, refactoring, documentation, and with accessibility guidance.

D) ChatGPT and comparable assistants may only be used in development for wording commit messages.

Question 27 (K2): Two implementations of the same search dialog look identical. Version 1 shows a loading state during an ongoing search, marks input errors at the affected field, has meaningful labels and a comprehensible focus order. Version 2 reacts to input errors without comment, has generic labels, and loses focus after submission. What does the comparison show about usability quality in implementation?

A) Both versions are equivalent, because usability quality is measured by the visual appearance.

B) Version 2 is preferable, because fewer system reactions make the interface appear calmer and tidier.

C) The difference is a pure testing topic and of no significance for implementation.

D) Usability quality arises essentially below the visible surface: System states, component logic, error messages, labels, focus order, and expectable behavior distinguish the two versions, although they look the same.

Question 28 (K2): A team introduces design tokens: Colors, spacing, and font sizes are defined as named values and referenced both in the design tool and in the code. Which description captures the function of this approach in the interplay with the design system correctly?

A) Design tokens are an export format for screenshots with which drafts are handed over to development without loss.

B) Design tokens make the design system superfluous, because the values already contain all design decisions.

C) Design tokens translate design decisions into consistent, technically usable values and thereby connect design and code; the design system additionally provides the rules, components, and patterns in which these values are used.

D) Design tokens are placeholder texts used in mockups until the final content is available.

Question 29 (K1): In a review guide, the accessibility basics for implementation are to be listed. Which list renders the points named in the learning objective?

A) WCAG conformity declaration, accessibility audit, user test with screen reader users, legal review, and certification.

B) Semantic HTML, sufficient contrast, meaningful labels, visible focus, keyboard operability, and ARIA only where necessary.

C) Dark mode, font size switcher, language selection, offline mode, and data minimization.

D) ARIA attributes on as many elements as possible, so that assistive technologies receive maximum information.

Question 30 (K3): For an AI-generated function “change payment details” in the customer account, the review procedure is to be defined. Which derivation of the review measures is appropriate?

A) Since the AI relies on proven libraries, linter runs and a brief look at the interface suffice.

B) A penetration test replaces all other reviews, because with payment data only security counts.

C) Code review by a second person, functional tests including the error cases (rejected change, session expiry, invalid inputs), UI tests, a security review because of the sensitive data, an accessibility check of the form, and professional acceptance against the acceptance criteria.

D) Approval after a successful run of the happy path in the test environment; the error cases are caught up on in operation.

Question 31 (K1): In a team discussion, various verification procedures are being mixed up. Which group of three contains exclusively fundamental evaluation methods for usability or expert-based procedures?

- A) User interview, online survey, and focus group.
- B) Load test, smoke test, and regression test.
- C) Heuristic evaluation, cognitive walkthrough, and usability test.
- D) Sprint retrospective, backlog refinement, and daily Scrum.

Question 32 (K2): After the implementation of a new document approval, a stakeholder asks: “The function is implemented and all tests are green, why do you now still want to evaluate and validate?” Which answer explains the purpose of evaluation and validation correctly?

- A) Evaluation and validation are formal mandatory steps of the exam regulations that must be documented before every release.
- B) Evaluation and validation repeat the functional tests with other tools in order to increase test coverage.
- C) Evaluation and validation serve to demonstrate the quality of the implementation to the stakeholder and to accelerate acceptance.
- D) Green tests show that the function runs as specified; evaluation and validation additionally check whether the solution is understandable and usable for the users, whether the assumptions behind the specification hold, and whether the goal is achieved in the real context of use.

Question 33 (K2): A developer is to organize a simple usability test for the new time-tracking screen for the first time and asks what she has to define for it. Which list names the central components completely?

- A) A research question, suitable test participants, a suitable test object with a defined system state, realistic test tasks, facilitation, observation, and the analysis of the results.
- B) A test lab with eye tracking, at least thirty test participants for statistical significance, and a video release from the legal department.
- C) A presentation of the screen, a feedback form with school grades, and a raffle among the participants.
- D) A completed requirements specification, the final version of the software, and the approval of the steering committee.

Question 34 (K2): Before a usability test, the facilitator considers whether she should ask the test participants to think aloud. Which assessment of the purpose and limits of think-aloud is correct?

- A) Thinking aloud should be avoided, because it distorts the measurement of task time and thereby devalues the entire test.
- B) Thinking aloud makes expectations, search strategies, and misunderstandings audible that pure observation does not show; at the same time, it can influence behavior, slow users down, and increase cognitive load, which is why it is used as a supporting method alongside behavioral observation.
- C) Thinking aloud replaces observation, because the test participants' statements fully explain their behavior.

D) Thinking aloud serves above all so that the facilitator can intervene immediately in case of problems and explain the right way.

Question 35 (K3): For the usability test of a health insurance app, it is to be checked whether insured persons can submit doctor's invoices independently. Which test task is worded in a goal-oriented way without prescribing the solution path?

- A) "Open the menu, select 'Benefits', and then tap the camera symbol to photograph the invoice."
- B) "Please test whether the submission function under 'Benefits' works correctly."
- C) "Find the menu item for invoice submission and describe how it is designed."
- D) "You were at the dentist yesterday and received an invoice for 180 euros. Use the app to do everything necessary so that the insurance can reimburse you the amount."

Question 36 (K3): A test of the export function yields three findings: (1) In a rare special case, records are silently omitted during export without users noticing. (2) A tooltip is slightly misplaced for many users, which irritates some. (3) An assistant step is misunderstood by about half, but can be corrected with one click. How does the team prioritize properly?

- A) Finding 1 first, although it is rare: Unnoticed incomplete exports can lead to wrong decisions based on the data; then finding 3 because of the high frequency of the misunderstanding; the tooltip last, since low impact despite frequency.
- B) Finding 2 first, because it affects the most users; frequency is the most objective prioritization criterion.
- C) Finding 3 first, because it is fastest to fix and the team thereby shows immediate visible progress.
- D) Adopt all three findings into the backlog with equal rank, because weighting test results endangers the objectivity of the analysis.

Question 37 (K1): After the release of a customer portal, the team wants to collect usability insights systematically. Which list names the typical sources specified in the learning objective?

- A) Competitor comparisons, analyst reports, trade fair contacts, press coverage, and industry studies.
- B) Commit history, code reviews, sprint reports, architecture decisions, and team retrospectives.
- C) Advertising impact analyses, click prices, newsletter open rates, social media likes, and campaign reach.
- D) User feedback, customer feedback, support tickets, usage data, technical signals, and UX metrics.

Question 38 (K2): A team wants to organize UX goals and suitable UX metrics for its learning platform and comes across the HEART framework. Which explanation of the basic idea is correct?

- A) HEART is a prioritization scheme that sorts backlog items into five urgency levels from H as in high to T as in trivial.
- B) HEART measures exclusively satisfaction via standardized surveys; the remaining dimensions are computed values derived from that.
- C) HEART structures UX goals and UX metrics along the dimensions happiness, engagement, adoption, retention, and task success; the team selects the relevant dimensions per product and derives suitable goals and measures for them, for the learning platform for example task

success in completing learning units and retention via the learners' return.

D) HEART is an approval process that defines which metrics a release must at least achieve before it may be shipped.

Question 39 (K2): The analysis shows: Users who watch the introduction video in full later use considerably more functions of the application. A stakeholder concludes: "The video makes the users more competent, we should make it mandatory." Which classification is correct?

A) The conclusion is correct, because the relationship rests on objective usage data and not on opinions.

B) The conclusion is correct, provided the relationship remains stable across two consecutive quarters.

C) The observation is meaningless, because usage data has no informative value without an accompanying usability test.

D) The observation is initially a pattern or correlation: Possibly the video works, but possibly motivated or experienced users are simply more likely to watch it in full. Cause, the resulting priority, and a reliable product decision such as an obligation require further examination.

Question 40 (K3): The usage data shows that the filter function of the results list introduced two months ago is barely used; occasional feedback calls it "well hidden". Whether demand is lacking or the function is not being found is unclear. How does the team properly carry this insight into further work?

A) It removes the filter function in the next release, because low usage proves lacking demand.

B) It formulates a discovery question with hypotheses (do users find the entry point? does filter demand even exist in typical tasks?), checks them for example through a small observation or targeted analysis, and depending on the result derives a concrete backlog item on findability or a conscious decision about the function.

C) It moves the most visible filter button to the top of the page and waits to see whether the metric rises; further clarification is only necessary if nothing changes.

D) It documents the low usage as a known state and waits until more users complain.

Part 2: Answer Key and Comments

For each question, the answer key names the associated learning objective and K level as well as comments on all four options.

Question 1 (LO 1.1, K1): Correct answer: C

A) Incorrect. Speed, freedom from errors, and style guide compliance are quality aspects of their own, but not a correct rendering of the three terms.

B) Incorrect. Usability is not limited to specific input devices, and user experience applies to all interactive systems, not only mobile applications.

C) Correct. This mapping corresponds to the syllabus's understanding of the terms: usability as an assessable quality of use, user experience as the more comprehensive usage experience, UX quality as the quality of use and usage experience in the development context.

D) Incorrect. Usability has not been superseded; it is part of user experience. UX quality is not a statement about visual modernity.

Question 2 (LO 1.3, K2): Correct answer: D

- A) Incorrect. The order of the events is not the problem; what is decisive is whether usability content is deliberately brought in and verified.
- B) Incorrect. Retrospectives improve collaboration in the team, but they do not replace substantive usability work on requirements, drafts, and increments.
- C) Incorrect. A strict Definition of Done does not prevent user orientation; on the contrary, it can even anchor usability criteria when these are included.
- D) Correct. That is exactly the core statement of the syllabus: The process framework guarantees no good UX; usability content must be actively brought into discovery, refinement, implementation, and verification.

Question 3 (LO 1.5, K2): Correct answer: D

- A) Incorrect. Better prompts can help, but they do not compensate for missing domain and context knowledge. The difference lies in pattern knowledge versus context knowledge.
- B) Incorrect. This is not about technical simplicity, but about whether a familiar, frequently occurring pattern is present or specific context of use is required.
- C) Incorrect. Economic relevance is not the reason; what is decisive is the availability of familiar patterns as opposed to unfamiliar domain context.
- D) Correct. That is exactly what the basic rule describes: strength with familiar patterns, weakness with unfamiliar context, where human context input and review are necessary.

Question 4 (LO 2.1, K1): Correct answer: D

- A) Incorrect. These terms come rather from product strategy and discovery tools, not from the chapter's basic vocabulary.
- B) Incorrect. Those are technical performance terms; they can influence use, but they do not belong to the required basic terms.
- C) Incorrect. These terms come from Scrum and agile work organization, not from the basic usability vocabulary.
- D) Correct. Exactly these terms are named in the learning objective as central basic usability terms.

Question 5 (LO 2.3, K2): Correct answer: A

- A) Correct. That is exactly what Chapter 2 describes: Perception is not complete and neutral, but steered by task, expectation, and situation; visibility alone is not enough.
- B) Incorrect. Red does not have a calming effect; the problem lies not primarily in the color choice, but in the focus of attention and the work situation.
- C) Incorrect. The interpretation as a discipline problem fails to recognize that selective perception is a normal human factor that design must take into account.
- D) Incorrect. Display size can play a role, but it does not explain why a large, central element is overlooked during focused activity.

Question 6 (LO 2.4, K2): Correct answer: B

- A) Incorrect. Click count can be an efficiency aspect, but it does not explain the interpretation and expectation problem of the label.
- B) Correct. The example shows how unclear terms work against mental models, violate conformity with user expectations, and increase cognitive load.
- C) Incorrect. Recognizability as a control is a different topic; here it is about the meaning of the action and the expected system behavior.

D) Incorrect. Experienced users also remain uncertain when a system is ambiguously labeled; moreover, the uncertainty causes real errors until then.

Question 7 (LO 2.9, K2): Correct answer: A

A) Correct. Predictions about user understanding are hypotheses without real verification, even when they are worded as statements of fact.

B) Incorrect. Training data about similar dashboards says nothing reliable about the specific users and their context.

C) Incorrect. An AI prediction is no user feedback; feedback comes from real people from real or anticipated use.

D) Incorrect. Precisely positive statements can lull into false security; open assumptions should remain visible.

Question 8 (LO 2.10, K3): Correct answer: D

A) Incorrect. Correct terminology and elaborate wording are no proof of substantive correctness.

B) Incorrect. Repeating the same answer does not increase reliability; the AI still has the same missing context.

C) Incorrect. A single error does not justify a general renunciation; what matters is the critical, checking approach.

D) Correct. Blanket judgments without context are a typical pattern of plausible-looking but incomplete AI answers; targeted re-checking and concrete, justified questions are the right approach.

Question 9 (LO 3.1, K1): Correct answer: A

A) Correct. Exactly these six terms are named in the learning objective as central discovery terms.

B) Incorrect. Those are work organization units from backlog tools, not discovery terms in the sense of the chapter.

C) Incorrect. Those are usage and web metrics; they can serve as a data source in discovery, but they are not the required basic terms.

D) Incorrect. This series mixes roles around purchasing and commissioning; it does not correspond to the six terms of the learning objective.

Question 10 (LO 3.2, K2): Correct answer: C

A) Incorrect. Architecture comes later; the upstream question is whether the solution addresses a real user problem.

B) Incorrect. This is not about rejection or market analyses, but about problem understanding before committing to a solution.

C) Correct. That is exactly the core of LO 3.2: Usability work begins before the solution idea with user, task, and context, otherwise a solution risks missing the need.

D) Incorrect. No such formal rule exists; what is decisive is the substantive argument, not a documentation obligation.

Question 11 (LO 3.4, K2): Correct answer: B

A) Incorrect. AI cannot reliably judge the credibility and representativeness of sources; that remains professional assessment.

B) Correct. Structuring, summarizing, making topics visible, and pre-formulating hypotheses is exactly the support described in LO 3.4.

C) Incorrect. Filling gaps with plausible assumptions is a risk, not processing; missing information must remain visible as missing.

D) Incorrect. Assessing customer value is a product and business decision and not the task of AI processing.

Question 12 (LO 3.6, K3): Correct answer: A

A) Correct. The wording separates cleanly: The abandonment is observed, the cause is a hypothesis with a defined verification path.

B) Incorrect. Analytics shows the what (abandonments), not the why; the causal claim is not proven.

C) Incorrect. The absence of counter-evidence does not turn an AI supposition into a finding; verification is required.

D) Incorrect. Causes are verifiable; precautionary changes without a hypothesis can miss the actual problem.

Question 13 (LO 4.1, K1): Correct answer: B

A) Incorrect. These artifacts belong rather to discovery and design, not to the refinement basic terms from LO 4.1.

B) Correct. Exactly these five terms are named in the learning objective.

C) Incorrect. Those are test management terms; acceptance criteria touch on tests, but the series does not correspond to the learning objective.

D) Incorrect. Those are Scrum terms for goal setting and work organization, not the chapter's refinement basic terms.

Question 14 (LO 4.2, K2): Correct answer: D

A) Incorrect. This is not about relief or brevity, but about making usage quality verifiable early.

B) Incorrect. Usability requirements are not directed against the business department; they complement the professional view with the usage perspective.

C) Incorrect. Subsequent improvements are permissible but often costly; the argument is professional and economic, not legal.

D) Correct. Usability is a quality requirement; if it is not anchored in the requirements, structural determinations arise that can later only be corrected at great expense.

Question 15 (LO 4.7, K2): Correct answer: C

A) Incorrect. The allocation is reversed: The user goal must be clarified before implementation, not only at the end; visual fine-tuning is no DoR topic.

B) Incorrect. Also reversed: Tests cannot be completed before implementation, and assumptions should be documented before the sprint begins.

C) Correct. Clarification before the sprint (DoR), verified quality at the end (DoD): That corresponds to the anchoring described in Chapter 4.

D) Incorrect. DoR and DoD are team agreements; precisely there, usability quality becomes visible and binding in the process.

Question 16 (LO 4.6, K3): Correct answer: A

A) Correct. Validation with an understandable error message, a visible system state after submission, and keyboard operability make central usability aspects verifiable.

B) Incorrect. Archiving and interfaces are technical or organizational criteria without a direct usage relation.

C) Incorrect. Processing time and audit log are performance and compliance criteria, not usability aspects in the sense of the LO.

D) Incorrect. Corporate design conformity is a design standard, but it addresses neither feedback nor error handling nor operability.

Question 17 (LO 4.8, K3): Correct answer: B

A) Incorrect. Concreteness is no proof of correctness; at least one criterion is professionally problematic and must not be adopted.

B) Correct. That is exactly what LO 4.8 demands: adopt, adapt, or exclude with justification, here with a clear exclusion reason (disclosure of account existence).

C) Incorrect. A blanket prohibition is unnecessary; AI suggestions are usable as a reviewed draft even for security-related functions.

D) Incorrect. A recognized problem must not be deliberately postponed until after the release; the assessment belongs in the refinement.

Question 18 (LO 5.1, K1): Correct answer: C

A) Incorrect. Those are discovery, market, and brand activities, not usability tasks in sprint implementation.

B) Incorrect. Those are development and operations tasks; they secure technical quality, not specifically the quality of use.

C) Correct. Exactly these tasks are named in the learning objective as typical usability tasks in implementation.

D) Incorrect. Those are organizational and personnel tasks without a usability relation.

Question 19 (LO 5.3, K2): Correct answer: D

A) Incorrect. Understandability is not the exclusive topic of a specialist role; testers check behavior and understandability alongside correctness.

B) Incorrect. The Scrum Master promotes working methods and visibility; the professional responsibility for UX quality does not therefore rest with them alone.

C) Incorrect. External support can complement, but the ongoing responsibility for usage quality remains in the team.

D) Correct. Exactly this division-of-labor description of the role contributions corresponds to Chapter 5: UX quality arises in the interplay, not in a single role.

Question 20 (LO 5.6, K2): Correct answer: A

A) Correct. Criteria orientation, structure, and recorded results make the difference from the taste discussion.

B) Incorrect. A design critique lives from multiple perspectives; it is not a format with speaking rights only for designers.

C) Incorrect. Anonymous writing is no defining feature; what is decisive is the reference to agreed criteria.

D) Incorrect. A design critique is an assessment and improvement format, not a majority vote on approvals.

Question 21 (LO 6.1, K1): Correct answer: B

A) Incorrect. These terms come from visual design and branding, but they do not cover the basic terms from LO 6.1.

B) Correct. Exactly these eight terms are named in the learning objective as basic terms of UI

design.

C) Incorrect. Those are performance terms; they influence use, but they are no UI design basic terms.

D) Incorrect. Those are accessibility terms; they are important, but they do not correspond to the required series of terms from LO 6.1.

Question 22 (LO 6.4, K1): Correct answer: C

A) Incorrect. The mappings are reversed: Low-fidelity serves early discussion, horizontal serves breadth, vertical serves depth.

B) Incorrect. This mapping mixes prototype types with technical and marketing questions.

C) Correct. This mapping corresponds to the prototype types named in LO 6.4 and their purposes.

D) Incorrect. Wireframe, mockup, and clickable prototype differ in level of detail and purpose; they are not synonyms.

Question 23 (LO 6.7, K2): Correct answer: D

A) Incorrect. AI masters many familiar patterns, not only search masks; what is decisive is the familiarity of the pattern, not the UI type as such.

B) Incorrect. AI is helpful with numerous standard patterns, such as login, form, dashboard, navigation, or onboarding.

C) Incorrect. The difference is systematic: familiar patterns versus unfamiliar domain context, not chance.

D) Correct. That is exactly what LO 6.7 explains: strength with frequent, well-documented patterns, weakness with domain-specific flows without context.

Question 24 (LO 6.6, K3): Correct answer: A

A) Correct. The assessment connects the criteria of the LO with the concrete context of use: public terminal, heterogeneous user group, high demands on perceivability, consistency, and accessibility.

B) Incorrect. A modern look is no review criterion when perceivability and understandability suffer.

C) Incorrect. Corporate design and technical completeness do not replace the check against the usability criteria.

D) Incorrect. A draft can and must be assessed against the criteria before operation; usage data comes in addition later.

Question 25 (LO 6.8, K3): Correct answer: B

A) Incorrect. Those are technical performance and compatibility questions, not the typical substantive weaknesses of AI-generated drafts.

B) Correct. States, permissions, error and abandonment paths, alternative usage situations, and accessibility are exactly the typical gaps named in LO 6.8.

C) Incorrect. Brand and quality impressions are no checkpoints for the typical professional weaknesses of AI drafts.

D) Incorrect. Tool and file hygiene is practical, but it does not address the substantive gaps of the draft.

Question 26 (LO 7.1, K1): Correct answer: C

- A) Incorrect. Selenium and Cypress are test frameworks, Docker and Kubernetes infrastructure tools; the assignment is inaccurate.
- B) Incorrect. Design tools deliver drafts; they are no AI coding tools, and automatic translation into production-ready code is not their core.
- C) Correct. Exactly this tool class and these fields of application are named in the learning objective for UI-related development.
- D) Incorrect. Assistance systems can support far more, such as code drafts, tests, documentation, and review questions; the syllabus describes no such restriction.

Question 27 (LO 7.3, K2): Correct answer: D

- A) Incorrect. The case shows precisely that the same appearance can mean very different usage quality.
- B) Incorrect. Missing feedback is not calm, but disorientation; users do not know what happened.
- C) Incorrect. These aspects must arise in implementation; tests can check them, but not create them afterwards.
- D) Correct. Exactly the quality aspects named in LO 7.3 make the difference between the visually identical versions.

Question 28 (LO 7.5, K2): Correct answer: C

- A) Incorrect. Tokens are no image handover, but named design values that are referenced in design and code alike.
- B) Incorrect. Values alone yield no design system; rules, components, states, and patterns come in addition.
- C) Correct. Exactly this bridging function between design, UI implementation, and code is described by LO 7.5.
- D) Incorrect. Placeholder texts are something else; tokens define design values such as colors, spacing, and typography.

Question 29 (LO 7.8, K1): Correct answer: B

- A) Incorrect. Audits, user tests, and conformity evidence are further-reaching measures, not the implementation basics from LO 7.8.
- B) Correct. Exactly these six points are named in the learning objective as accessibility basics in implementation.
- C) Incorrect. These functions can be useful, but they are no fundamental accessibility basics in the sense of the LO.
- D) Incorrect. ARIA should be used precisely only where native elements are not sufficient; excessive ARIA can worsen accessibility.

Question 30 (LO 7.10, K3): Correct answer: C

- A) Incorrect. Linter and visual inspection cover neither security nor usability nor professional risks.
- B) Incorrect. Security is central, but understandability, error handling, accessibility, and professional correctness must also be checked.
- C) Correct. The combination of code review, functional tests with error cases, UI tests, security review, accessibility check, and professional acceptance corresponds to LO 7.10, appropriate for a sensitive function.

D) Incorrect. With payment data, known error cases must not be pushed into productive operation.

Question 31 (LO 8.1, K1): Correct answer: C

A) Incorrect. Interviews, surveys, and focus groups are research and questioning methods; they deliver insights, but they are not the evaluation procedures named in the LO.

B) Incorrect. Those are technical test types for stability and correctness, not usability evaluation methods.

C) Correct. Exactly these three procedures are named in the learning objective: two expert-based procedures and the usability test with real or representative users.

D) Incorrect. Those are Scrum events for collaboration and improvement, not evaluation methods.

Question 32 (LO 8.2, K2): Correct answer: D

A) Incorrect. This is not about formal obligations, but about the substantive verification of usability and assumptions.

B) Incorrect. Evaluation and validation repeat no functional tests; they check a different dimension: use.

C) Incorrect. Demonstration and acceptance are not the professional purpose; the point is insight, not persuasion.

D) Correct. Implemented does not mean usable: Evaluation and validation check understandability, usability, and the viability of the assumptions in the context of use.

Question 33 (LO 8.6, K2): Correct answer: A

A) Correct. Exactly these components are named in the learning objective; they make even a simple, pragmatic test viable.

B) Incorrect. Lab, large samples, and special technology are not required for a simple usability test; even a few suitable test participants uncover important problems.

C) Incorrect. Presentation and grading capture opinions instead of observable behavior; that is no usability test.

D) Incorrect. A test needs neither final software nor committee approvals; it can take place early with prototypes and a defined state.

Question 34 (LO 8.8, K2): Correct answer: B

A) Incorrect. When task times are in the foreground, one can respond to that; a general renunciation, however, gives away the insights into thought processes.

B) Correct. Exactly this dual role is described by LO 8.8: valuable insights into mental models alongside known limits (influence, slowdown, load).

C) Incorrect. Saying and doing can diverge; observing actual behavior remains indispensable.

D) Incorrect. Facilitation is precisely not meant to lead to the solution; hesitating and searching are important observations.

Question 35 (LO 8.7, K3): Correct answer: D

A) Incorrect. The task prescribes every operating step; whether users find and understand the way themselves is not tested.

B) Incorrect. The task names function and menu item and turns the test participant into a function checker instead of a user with her own goal.

C) Incorrect. The task reveals that a menu item exists and asks about design instead of goal

achievement.

D) Correct. The task describes a realistic situation with a goal from the user's view, without anticipating terms or controls of the app; success is observable.

Question 36 (LO 8.10, K3): Correct answer: A

A) Correct. The weighing connects severity, risk, and user impact with frequency: A rare problem with possible consequential damage without visibility weighs more heavily than frequent but minor irritations.

B) Incorrect. Frequency alone is no sufficient criterion; a rare but consequential problem can take precedence.

C) Incorrect. Low effort is a legitimate factor, but it must not outvote severity and risk.

D) Incorrect. Prioritization is no violation of objectivity, but necessary professional assessment; without it, the clear next step is missing.

Question 37 (LO 9.1, K1): Correct answer: D

A) Incorrect. Market and outside views can be strategically interesting, but they are not the typical sources of usability insights from real use.

B) Incorrect. These sources document the development, not the use of the product after release.

C) Incorrect. Those are marketing metrics; they say little about the usage quality of the product.

D) Correct. Exactly these six sources are named in the learning objective for insights into usability after release.

Question 38 (LO 9.4, K2): Correct answer: C

A) Incorrect. HEART is no urgency or prioritization scheme; the letters stand for dimensions of the user experience, not for levels.

B) Incorrect. Happiness is only one of the five dimensions; the others rest on actual usage behavior, not on survey derivations.

C) Correct. That is exactly the basic idea: five dimensions as a structuring aid, from which product-specific relevant goals and metrics are derived.

D) Incorrect. HEART defines no release thresholds; it organizes goals and metrics, but it replaces no approval processes.

Question 39 (LO 9.6, K2): Correct answer: D

A) Incorrect. Objective data too shows initially only a relationship, no direction of effect.

B) Incorrect. Stability over time does not turn a correlation into a cause; the self-selection would remain.

C) Incorrect. Usage data is valuable; it must only be classified correctly and, where needed, complemented by further analysis.

D) Correct. Exactly this distinction is demanded by LO 9.6: separating pattern and correlation from cause, priority, and a reliable product decision.

Question 40 (LO 9.8, K3): Correct answer: B

A) Incorrect. Low usage can also be due to poor findability; removal would be a decision on an unclarified cause.

B) Correct. Carrying the insight into a discovery question and hypotheses, checking in a targeted way, then deriving concrete backlog items or decisions: exactly this path corresponds to LO 9.8.

C) Incorrect. A change without clarifying the cause is a guessing game; if the demand is lacking,

even the best placement will not help.

D) Incorrect. Waiting gives away the available hints; insights should be actively carried into questions, measures, or decisions.