

DeepQuali: Initial results of a study on the use of large language models for assessing the quality of user stories.

Dr. Adam Trendowicz

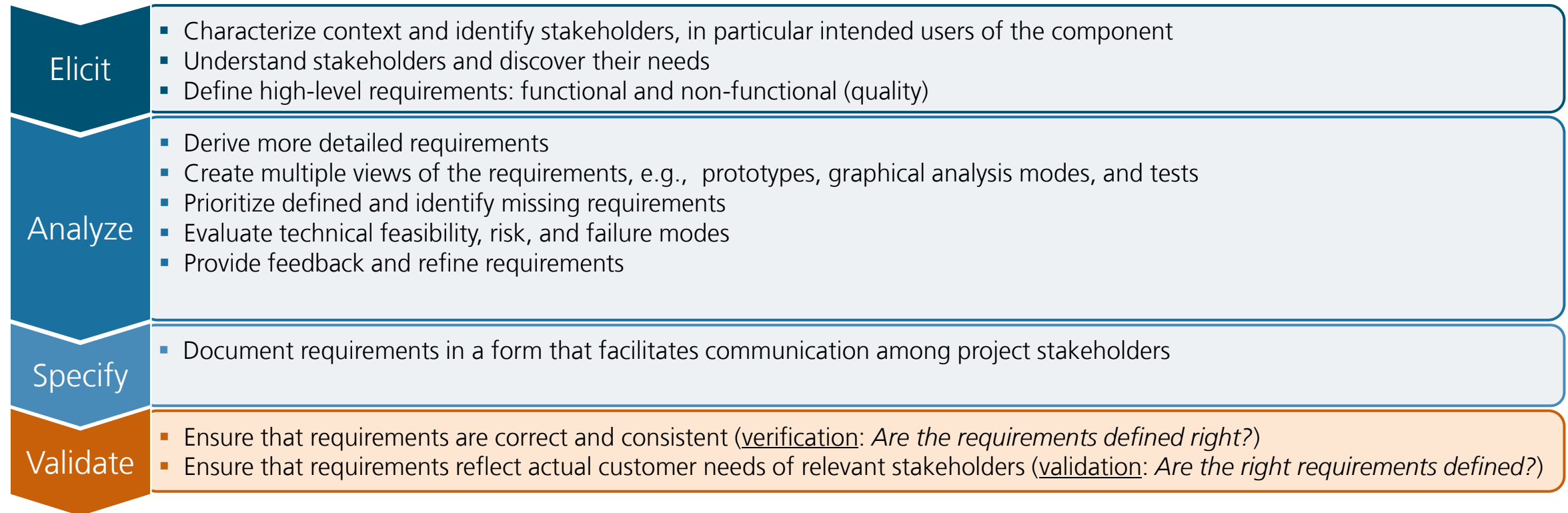
Dep. Data Science (DS)

Fraunhofer IESE

Workshop on Applications of Artificial Intelligence for Software Engineering (AI4SE)
49th German Conference on Artificial Intelligence, **KI2026**
Bremen, 12. August 2026

Focus of DeepQuali

DeepQuali focuses on the validation of software requirements (ISO/IEC 15288:2023)



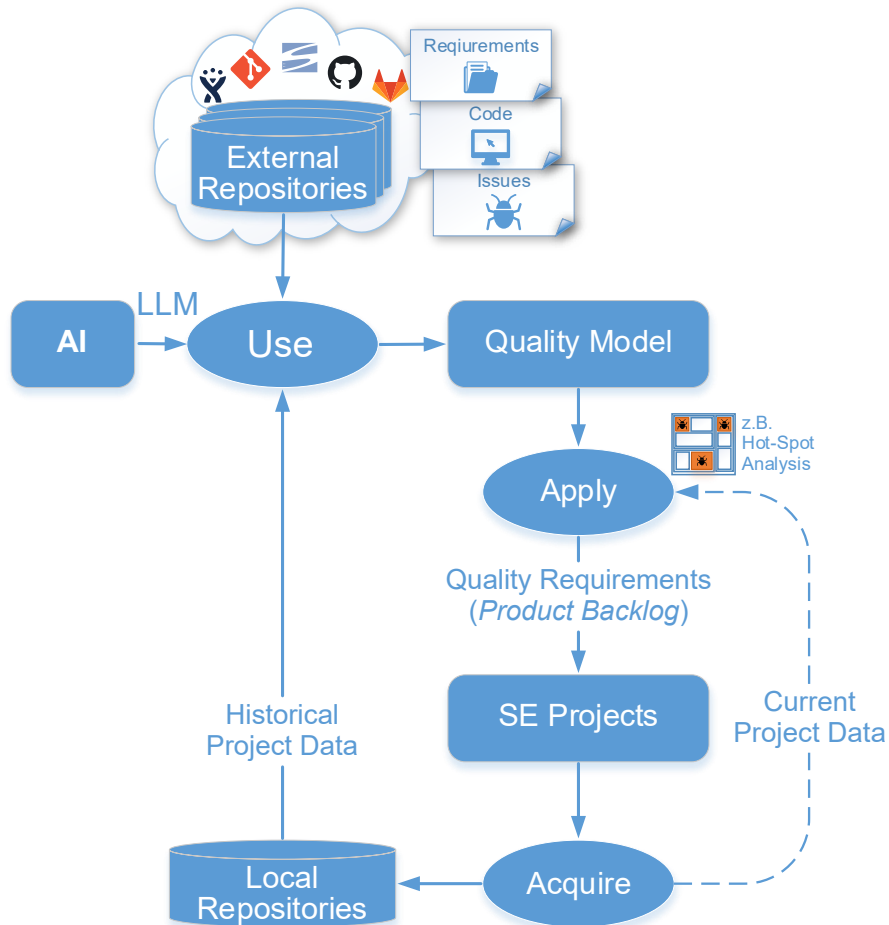
Requirements Validation Tasks Supported by AI

AI can be used for several specific evaluation tasks, which can be organized into an evaluation pipeline

Purpose/Task	AI Usage	Input	Output
Assess quality of requirements	Rate the quality of requirements	• Quality model, e.g., criteria, rules • Context, e.g., domain, infrastructure	Quality assessments, e.g., scores
Detect requirements quality deficits	Describe quality deficits of requirements	• Quality model, e.g., criteria, rules • Context, e.g., domain, infrastructure	Specific quality deficits (problems, issues)
Classify quality deficits, e.g., criticality	Group requirements deficits	• Quality deficits • Classification criteria	Classification of quality deficits
Derive problem resolutions	Generate potential resolutions of given quality deficits	• Quality assessments / deficits • Classification of quality deficits • Solution criteria	• Resolution recommendations: How requirements can be improved • Improved requirements: Examples & alternatives
Explain (reason) the results	Explain & justify specific outcomes of the quality evaluation, e.g., ratings, deficits, classification, or resolutions	Outputs of the above tasks: • Quality assessments • Quality deficits • Classification of quality deficits • Problem resolutions	Reasoning underlying a specific output

DeepQuali Projekt

AI-based assessment and improvement of software quality



Context

- Joint research & innovation project with two SMEs, funded by German Ministry of Research, Technology and Space.

Objective

- Development-time assessment and improvement of software quality, in the context of agile software development, based on the direct analysis of unstructured software artifacts, using AI technologies.

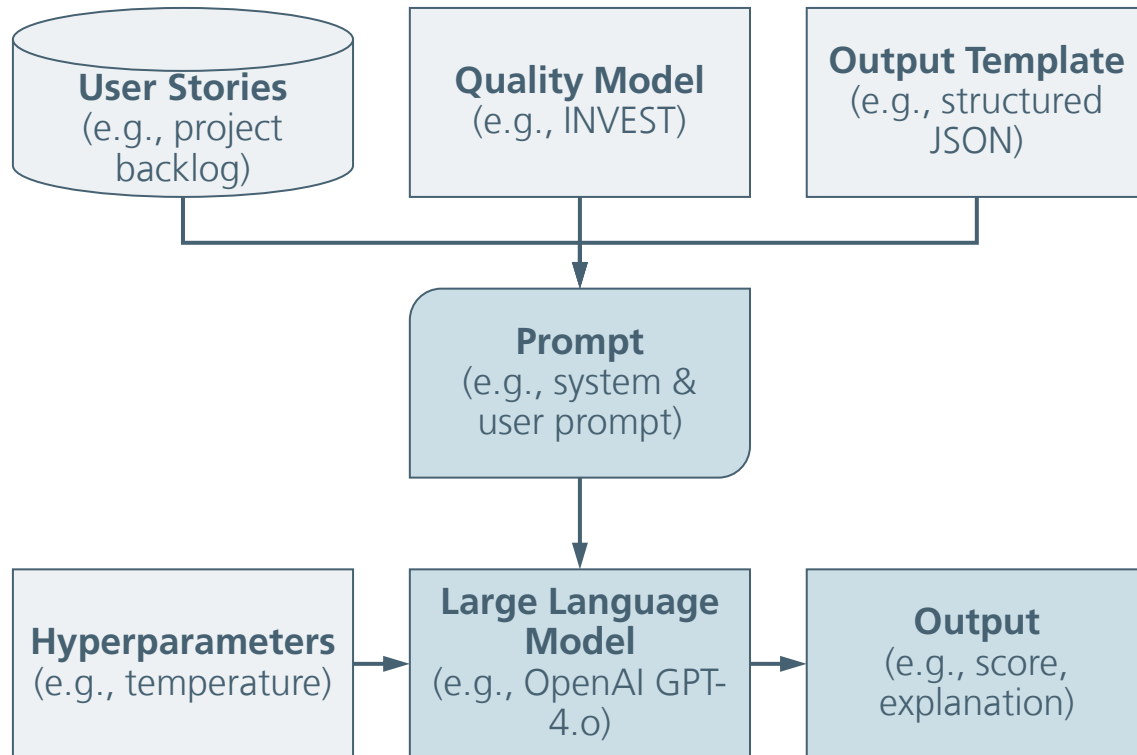
Components

- Requirements: Quality assessment of requirements specifications, specifically user stories.
- Source code: Quality assessment of software architecture and code, based on the analysis of source code.

Web: <https://deepquali.de/>

DeepQuali Approach

Using LLM upon custom-defined decision criteria to assess the quality of user stories



Context

- Small and medium enterprises
- Agile software development

Need

- Effective and efficient control of user stories' quality

Objectives

- Assess quality, identify deficits, propose improvements, explain reasoning
- Assure comprehensibility

Solution

- LLM: OpenAI GPT4.o
- Structured inputs & outputs (JSON)

DeepQuali Inputs and Outputs

Forcing structured form of inputs and output improves explainability of the approach

```
{
  "user_story": "",
  "evaluation": {
    "llm_parameters": {
      "datetime": "...",
      "model": "gpt-4o",
      "seed": 1337,
      "temperature": 0.0,
      "top_p": 1.0,
      "max_tokens": null,
      "stop": [],
      "presence_penalty": 0.0,
      "frequency_penalty": 0.0
    },
    "invest": {
      "evaluations": [
        {
          "name": "Name of the INVEST criterion.",
          "explanation": "Explanation of the quality assessment.",
          "score": "Quality assessment on the scale 1 (worst) to 4 (best).",
          "problems": [
            {
              "explanation": "Textual explanation of the problem.",
              "actionable_item": "Recommended improvements to the use story.",
              "severity": "Problem's severity on the scale 1 (lowest) to 5 (highest)."
            }
          ]
        }
      ]
    },
    "ready_to_implement": {
      "explanation": "Text explanation of the RTI assessment.",
      "score": "RTI rating on the scale 1 (lowest) to 4 (highest)."
    }
  }
}
```

The input user story, potentially structured into identifier, title, description, preconditions, acceptance criteria, story points estimate, etc.

Execution parameters of the LLM model used.

Assessment of the user story's quality on each INVEST criterion, including textual explanation, numerical score, specific problems, their severity, and recommended resolutions.

Name of the INVEST criterion.

Explanation of the quality assessment.

Quality assessment on the scale 1 (worst) to 4 (best).

Specific quality problems regarding the given criterion.

Textual explanation of the problem.

Recommended improvements to the use story.

Problem's severity on the scale 1 (lowest) to 5 (highest).

Assessment of the extent to which the user story is ready to be implemented, either using implicit definition of RTI or explicitly defined criteria (custom definition of ready, DOR).

Text explanation of the RTI assessment.

RTI rating on the scale 1 (lowest) to 4 (highest).

Inputs

- Parameters of the LLM model
- User story: anonymized, in original structure
- Assessment criteria: INVEST, DOR (Definition Of Ready)

Outputs

For each criterion:

- Score: Numeric quality rating of fulfilling given quality criterion
- Explanation: Explanation and justification of the numeric quality score
- Problems: description of specific quality issues regarding the criterion
- Severity: Numeric rating of problems' severity
- Actionable item: Concrete improvements recommended to solve problems

Based on all criteria:

- Ready to implement: readiness of a given user story to be released for implementation
 - Score: Numeric rating the readiness to be implemented
 - Explanation: Explanation and justification of the numeric quality score

Prompts Used in DeepQuali

Decomposing prompt into generic and specific part improves portability of the approach

System Prompt

System prompt represents the generic part of the prompt describing the context and the task

- **Role:** You are an experienced software engineer with deep knowledge in agile practices and user story evaluation.
- **Task:** You evaluate user stories on whether they are ready to be implemented.
- **Outcome:** You respond in a given JSON schema.

User Prompt

Specific part of the prompt explaining details of the task and its expected outcomes

- I will provide you with a user story below.
- Your task is to carefully read through the user story and evaluate it using the INVEST criteria.
- If an INVEST criterium is completely fulfilled and there is nothing to complain about, please don't hesitate to give the score 10.
- If an INVEST criterium is not completely fulfilled, please provide a brief explanation of what is missing and what can be done to improve it.
- Here is the user story in JSON: `{user_story_json}`

Used in the study command-line prototype was later implemented as GUI-Application

Used in the study command-line prototype was later implemented as GUI-Application

Projekte

T

Projekte

Teams

Projekt hinzufügen

Projekte > DeepQuali > User S

Projekte > DeepQuali > Epics

Projekte > DeepQuali > Global

DeepQuali

DeepQuali ☆

DeepQuali ☆

DeepQuali ☆

Übersicht

Übersicht

Code-Quali

Übersicht

Code-Qualität

Übersicht

Code-Qualität

Code-Architektur

User Stories

Epics

Global

Filtern nach

ID: US08 – US08. Projekt

ID: E1 – E1. Implementier

ID: DeepQuali 10

← 1 / 1 →

Suche

Suche nach

Sortierung

Name (A–Z)

Bewertung

10

9

8

7

6

5

4

3

2

1

Analyse der Findings

Kriterium

Beweis

Estimable 6

Erklärung

The story's broad scope and focus are well-defined.

Fix-Vorschlag

Break down the story into smaller, more specific user stories.

Independent 5

Erklärung

The story bundles multiple features or states that could be separated.

Fix-Vorschlag

Split the story into smaller, independent pieces.

Negotiable 8

Erklärung

Criterion not fully satisfied, but it's negotiable.

Fix-Vorschlag

Add concrete details and acceptance criteria.

Small 4

Erklärung

The story is too large and complex.

Fix-Vorschlag

Divide the story into smaller, simpler tasks.

Testable 6

Valuable 9

Analyse der Findings

Kriterium

Be

Zusammenhängend 6

Erklärung

Minor potential for improved clarity.

Fix-Vorschlag

Ensure user stories explicitly link related functionality.

Vollständig 7

Erklärung

User stories do not fully cover all aspects of the feature.

Fix-Vorschlag

Add user stories to cover projected gaps.

Nachvollziehbar 8

Erklärung

The epic description leaves some details unclear.

Fix-Vorschlag

Clarify user roles and edge cases.

Analyse der Findings

Kriterium

Bewertung

Erklärung

Konsistent 10

The user stories cover different aspects of the platform's functionality without conflicting requirements or descriptions. Each story complements the others in the user journey and feature set.

Nicht redundant 9

Most user stories address distinct features or states of the platform. However, some stories describe related or overlapping functionality, such as dashboard states (US02, US05), project overview states (US08, US10, US11, US12), and code quality findings (US13, US14, US15). These are complementary but could be reviewed for potential consolidation or clearer separation.

Erklärung

[Story US02](https://example.com/DQ-8) and [Story US05](https://example.com/DQ-11) both describe dashboard states (empty and populated) which are related but distinct. Similarly, [Story US08](https://example.com/DQ-14), [Story US10](https://example.com/DQ-16), [Story US11](https://example.com/DQ-17), and [Story US12](https://example.com/DQ-18) describe various project overview states that overlap in content and could be consolidated for clarity.

Fix-Vorschlag

Review and possibly consolidate dashboard and project overview user stories to reduce overlap and improve clarity.

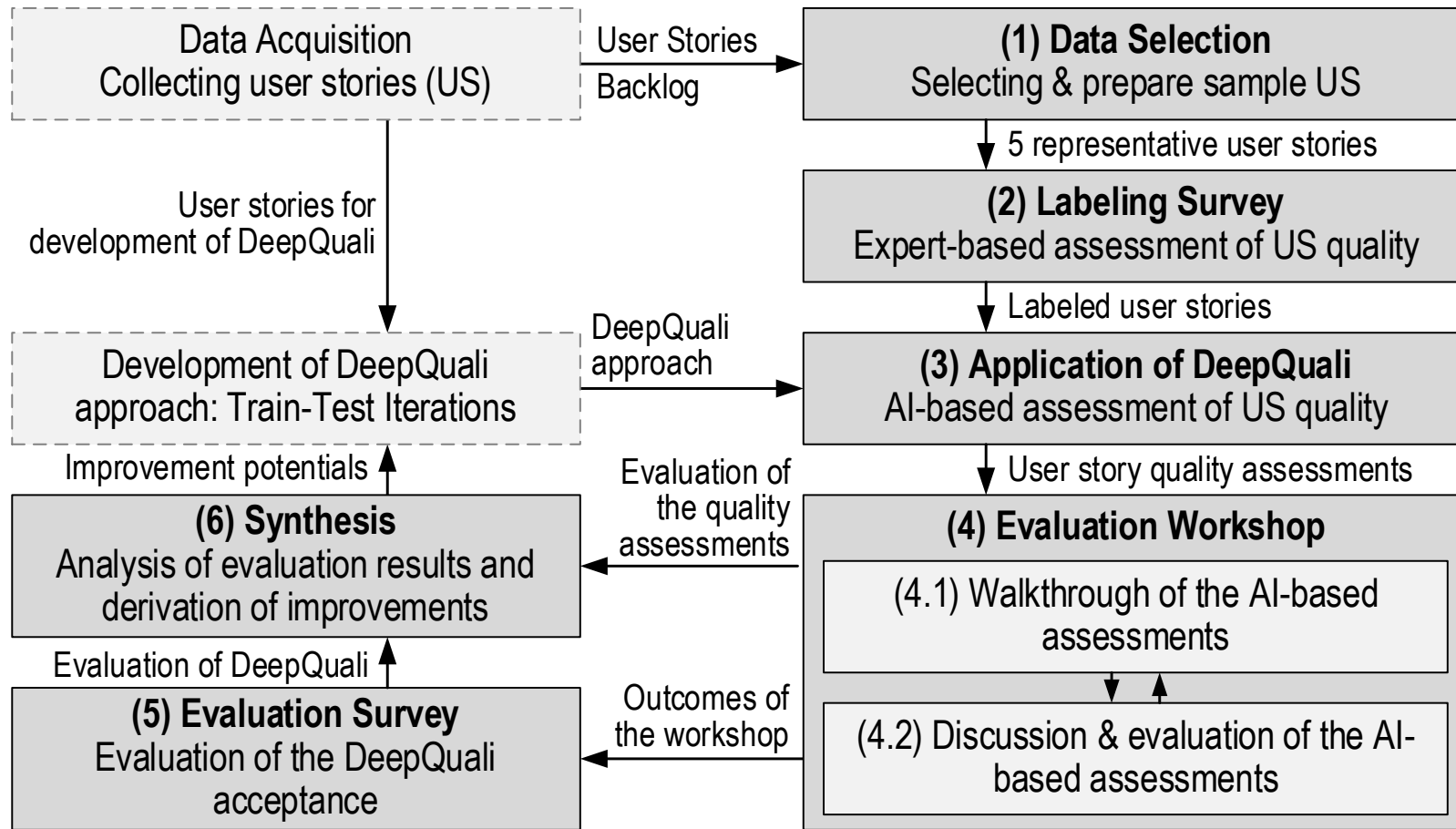
Informationen

Gefunden am
vor 3 Monaten

Aktueller INVEST-Score
100 %

DeepQuali Evaluation Procedure

Evaluation is based on the ground truth provided by subject matter experts

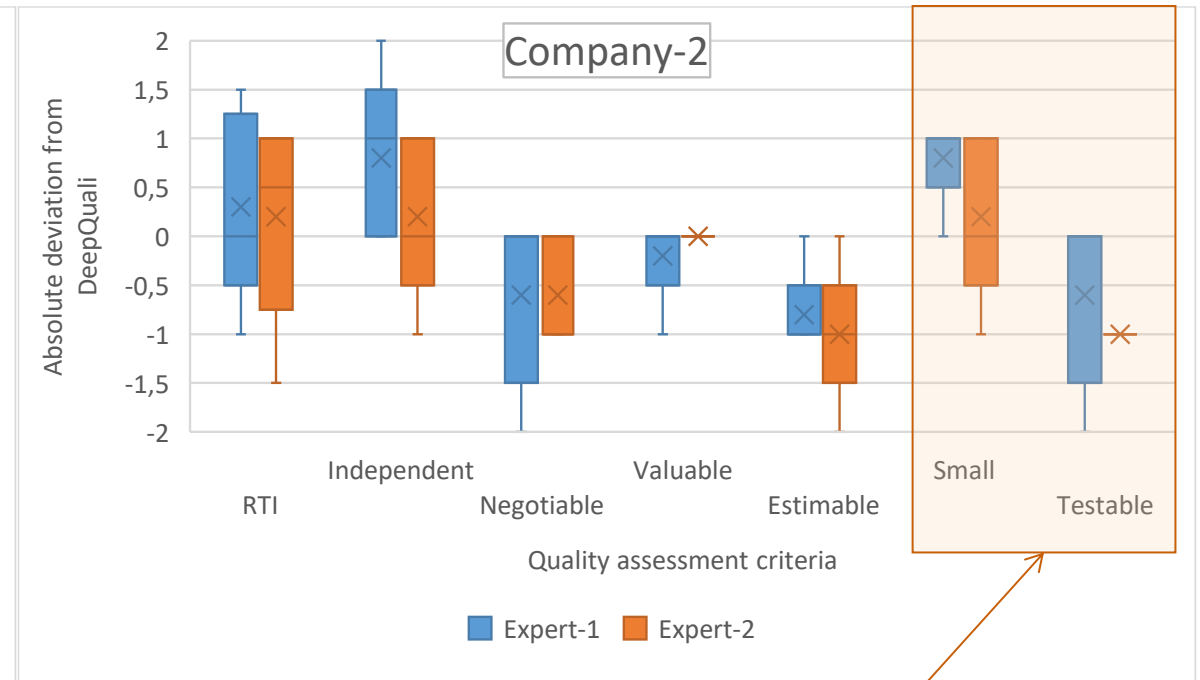
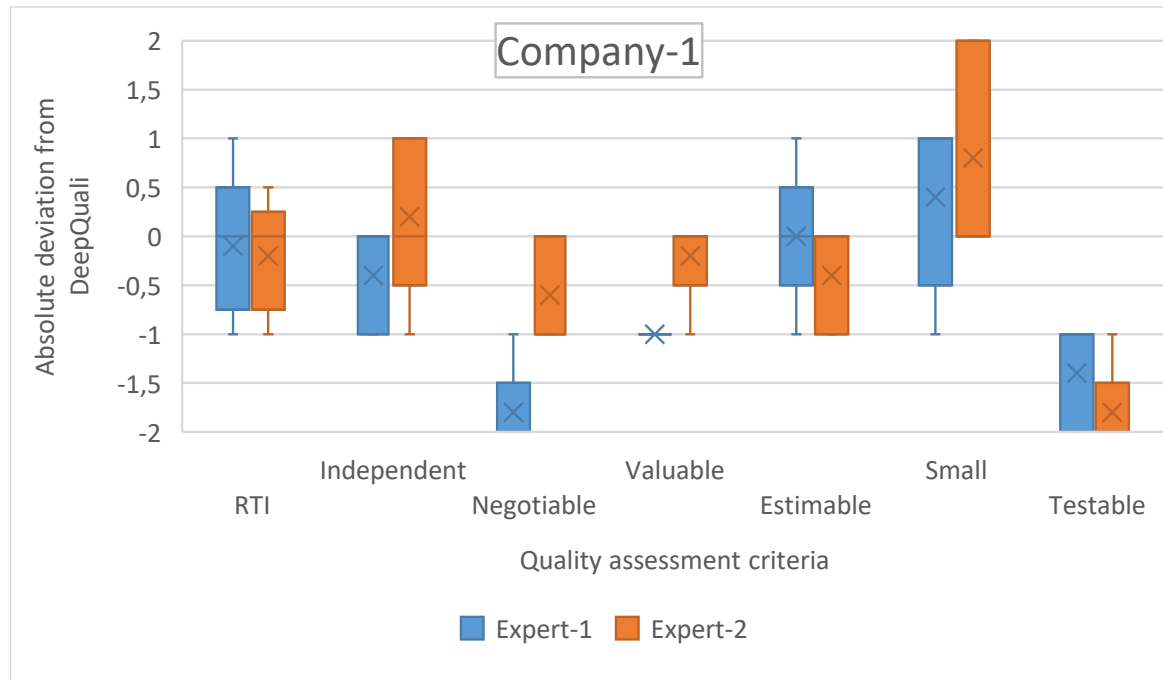


Evaluation context

- Two small-size enterprises (SMEs)
- Agile software development
- Four human experts

DeepQuali Evaluation Results: Consistency of the quality assessments

Experts are more critical on the *Testable* and less critical on the *Small* criteria

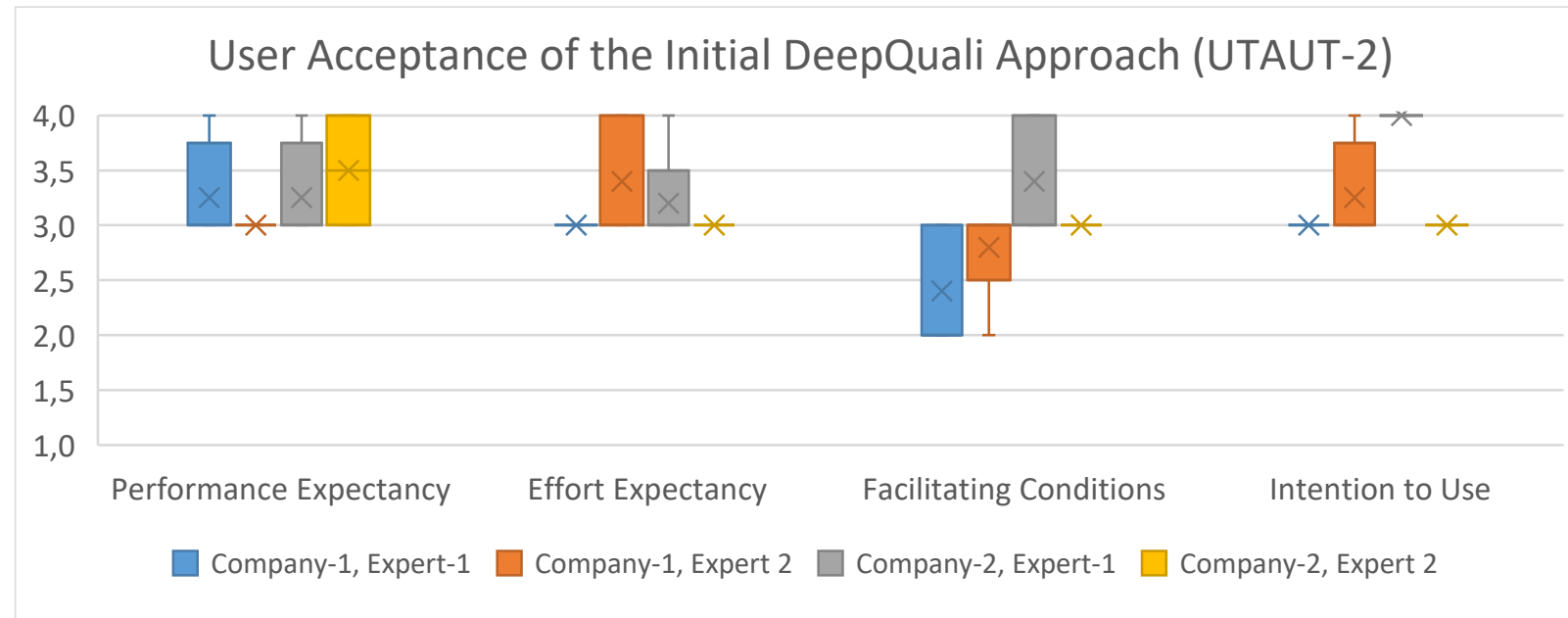


General observations

- Expert E1 with higher experience was principally more critical than AI
- Experts consistently more critical than AI on the *Testable* criterion
- Experts consistently less critical than AI on the *Small* criterion

DeepQuali Evaluation Results: Users' Acceptance

Experts appreciate the tool's practical usefulness but doubt its integrability into their work environment



Survey based on UTAUT

- **Performance expectancy:** The degree to which a person believes that using the technology will help improve their job or task performance.
- **Effort expectancy:** The perceived ease of using the technology.
- **Facilitating conditions:** The extent to which a person believes that organizational and technical resources exist to support technology use.
- **Intention to use:** The person's intention or willingness to use the technology.

UTAUT: Unified Theory of Acceptance and Use of Technology

DeepQuali Evaluation Results: Users' Feedback

Organization-specific standards may differ from 'average' baseline of an LLM

- High agreement between experts and LLM on INVEST criteria.
- Highest deviation in assessments on NET: *Negotiable*, *Estimable*, and *Testable* criteria.
 - Experts are more critical in their assessments than LLM.
 - More experienced experts are more critical than those with less experience.
- Experts rated consistently higher than LLM only on *Small*.
 - Potential reasons: The user stories represented a typical ("standard") size, the experts were familiar with. LLM expected smaller user stories.
- Human experts found the specific deficits identified by the LLM, along with additional explanations, very helpful for understanding the quality of user stories.

Lessons Learned

Proper framing, rather than brute force, improves the quality of LLM outcomes

- Provide LLM with **comprehensive context** information
 - Development and operation context of the software system, e.g., application domain and platform
 - Internal quality standards (e.g., regarding the *Size*), which may deviate from the common ones
- Provide relevant information to LLM as both **general specification and specific examples**.
 - Example: Specification of expected structure and size of user stories as well as typical examples of small-, medium- and large-size user stories.
- Force LLM to explain its **outcomes** and its **reasoning**
 - (!) It **improves the assessment outcomes** (quality rating closer to experts' assessments)
 - It **improves understanding** and **trustworthiness** of the outcomes

Thank you!

Questions, Remarks?





Contact

Dr. Adam Trendowicz
Fraunhofer IESE
Fraunhofer-Platz 1
67663 Kaiserslautern

DeepQuali Project Partner



LET'S DEV

Parts of this work were funded by the Federal Ministry of Research, Technology and Space (BMFTR) under grant no. 01IS23016D.