



Beyond Accuracy

Adaptive Decision Support Systems



Beyond accuracy

“Making recommendations better”

There are aspects beyond the algorithm that determine the quality of recommender systems

“Being accurate is not enough”

There are metrics beyond accuracy that should be used to build and evaluate recommender algorithms

“Explaining the user experience of recommender systems”

A framework to structure the evaluation of recommenders

A man with a shaved head, wearing a white shirt and a dark jacket, is speaking at a podium. The podium has a logo that reads "NYU STERN" and "NEW YORK UNIVERSITY - LEONARD N. STERN SCHOOL OF BUSINESS". The background is a light-colored wall with a grid pattern.

“The algorithm accounts for only 5% of the commercial success of our recommender systems [...] The interactive components of a recommender account for about 50%”



HRI

Three pillars of Human Recommender Interaction (HRI)

The act of getting recommendations (dialogue)

Perception of the recommender over time (personality)

How the recommender meets the user's needs



Recommendations

Recommendations shouldn't just be **correct**, they should be **notable** (cf. serendipity) and **useful**

In system terms, this might require the recommender to not just be **personalized**, but to be **bold** and take **risks**

Affirmation is good, but prevent **pigeon-holing**

The system should **adapt** to changing needs and provide **fresh** recommendations



Process

The recommendation process must be **transparent** and **usable**

Inspire **trust** in users

Meet users' **needs** and **expectations** w.r.t. their task



Contributions

What does this paper contribute to the state-of-the-art?

Mindset: think about recommenders beyond the algorithm

Agenda-setting: outline goals and objectives of a recommender system (and research)

What else?



Criticisms

What is missing from this paper?

Practical design advice: how do we build a recommender to accomplish these goals?

Evaluation advice: How do we evaluate these qualities?

What else?



NBC

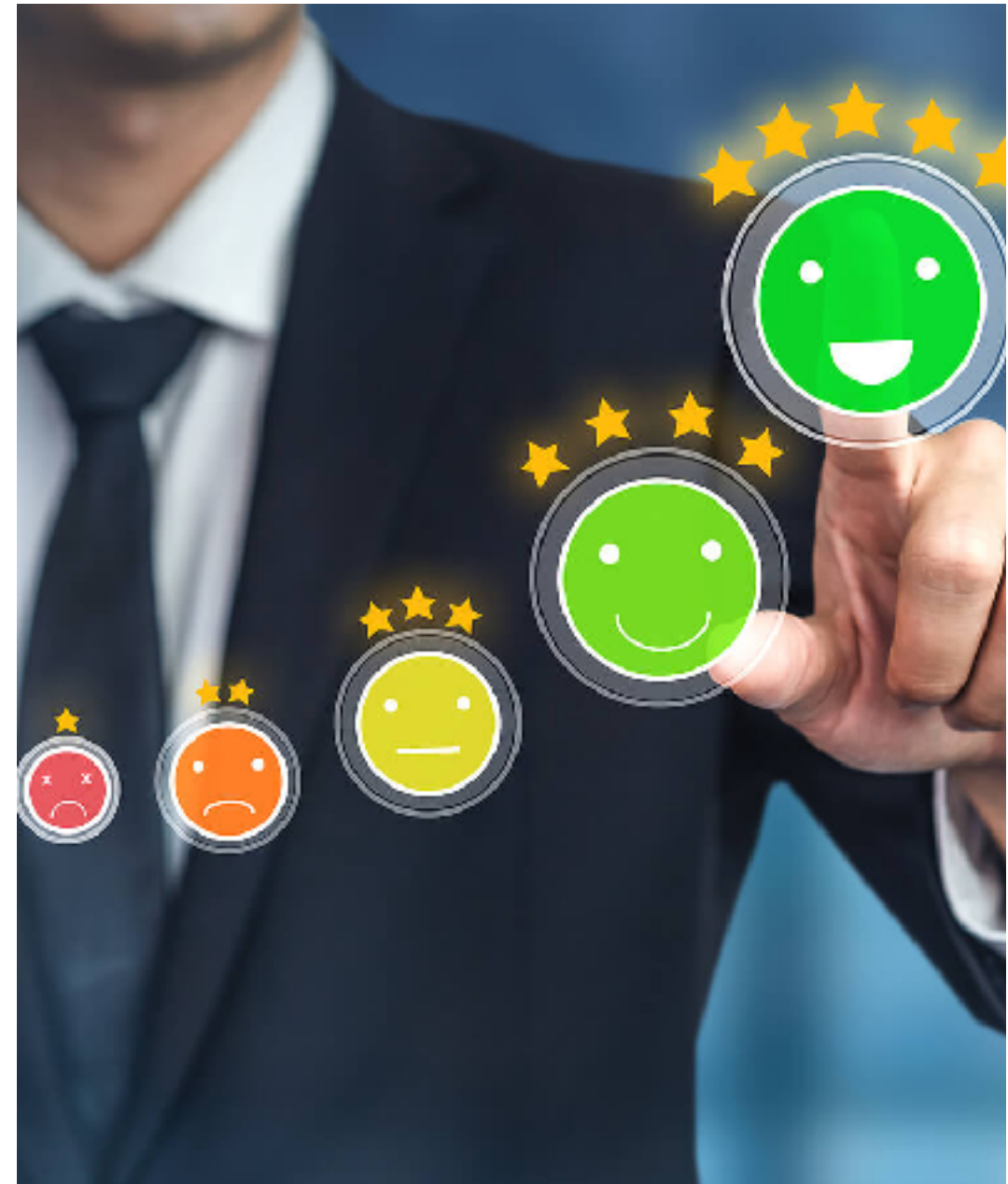
Left Field



Accuracy?

McNee et al. (2002): Item-Item CF provided the best predictions, but was rated least helpful by users!

Torres et al. (2004): CBF-separated CF had the lowest predictive accuracy, but resulted in the highest user satisfaction!





Beyond accuracy

Recognizes that accuracy is not the only way to evaluate a recommender

- Does not consider list composition

- Does not consider serendipity

- Does not consider user experiences and expectations



List composition

The recommendations are based on similarity

Remember w_{ij} and w_{uv} ?

But they should also be diverse!

Compared to the rated items (non-obvious recommendations)

Compared to the other items in the list (this also reduces choice overload)

Over time (adaptability and freshness)

To some extent, diversity is the opposite of similarity!



List composition

If we consider diversity the opposite of similarity, you simply get the opposite of accuracy...

How can we reconcile diversity and similarity?

Use different metrics: e.g. user-user or item-item similarity
+ genre or emotion diversity

Diversify perpendicular to the user vector (see next slides)



List composition

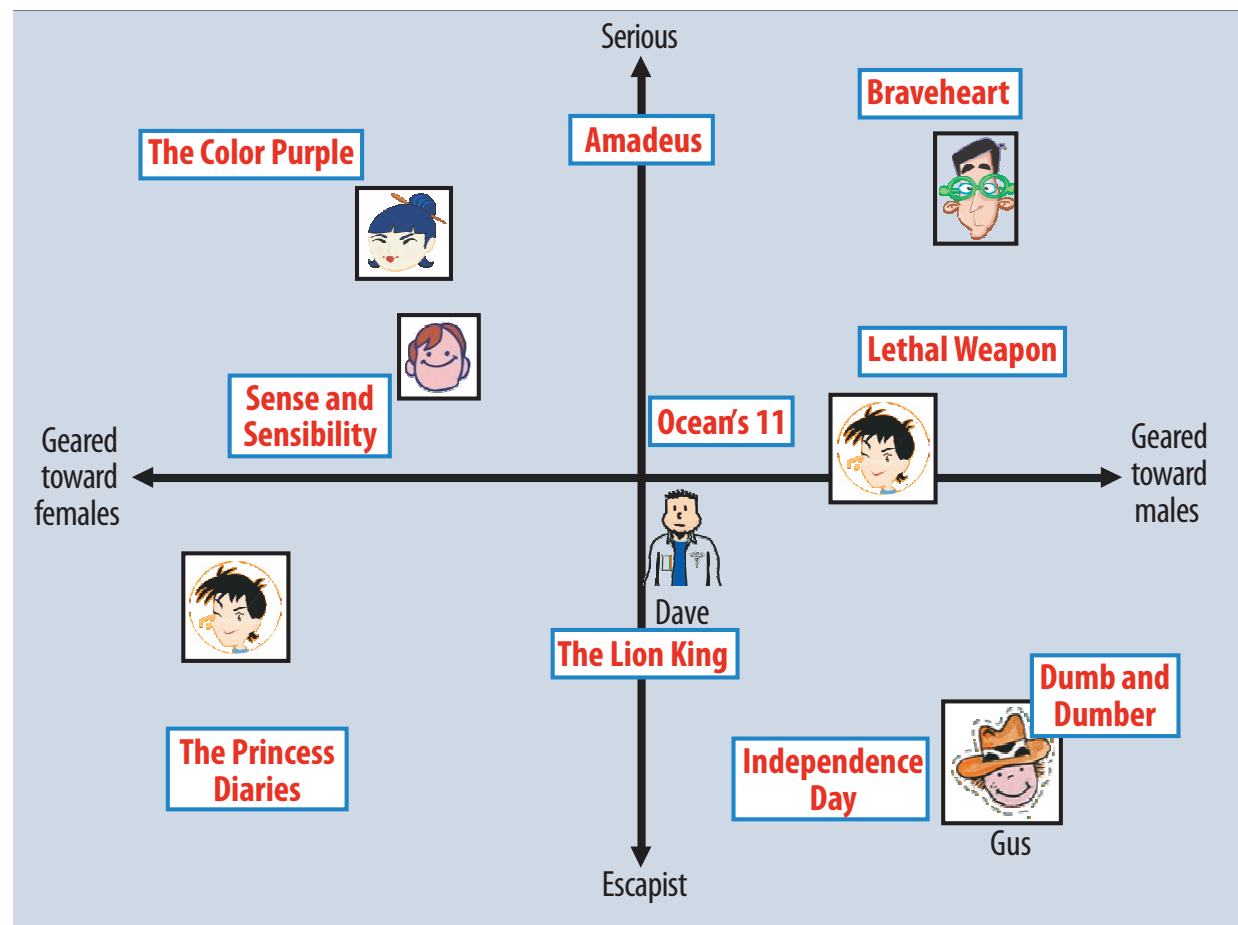
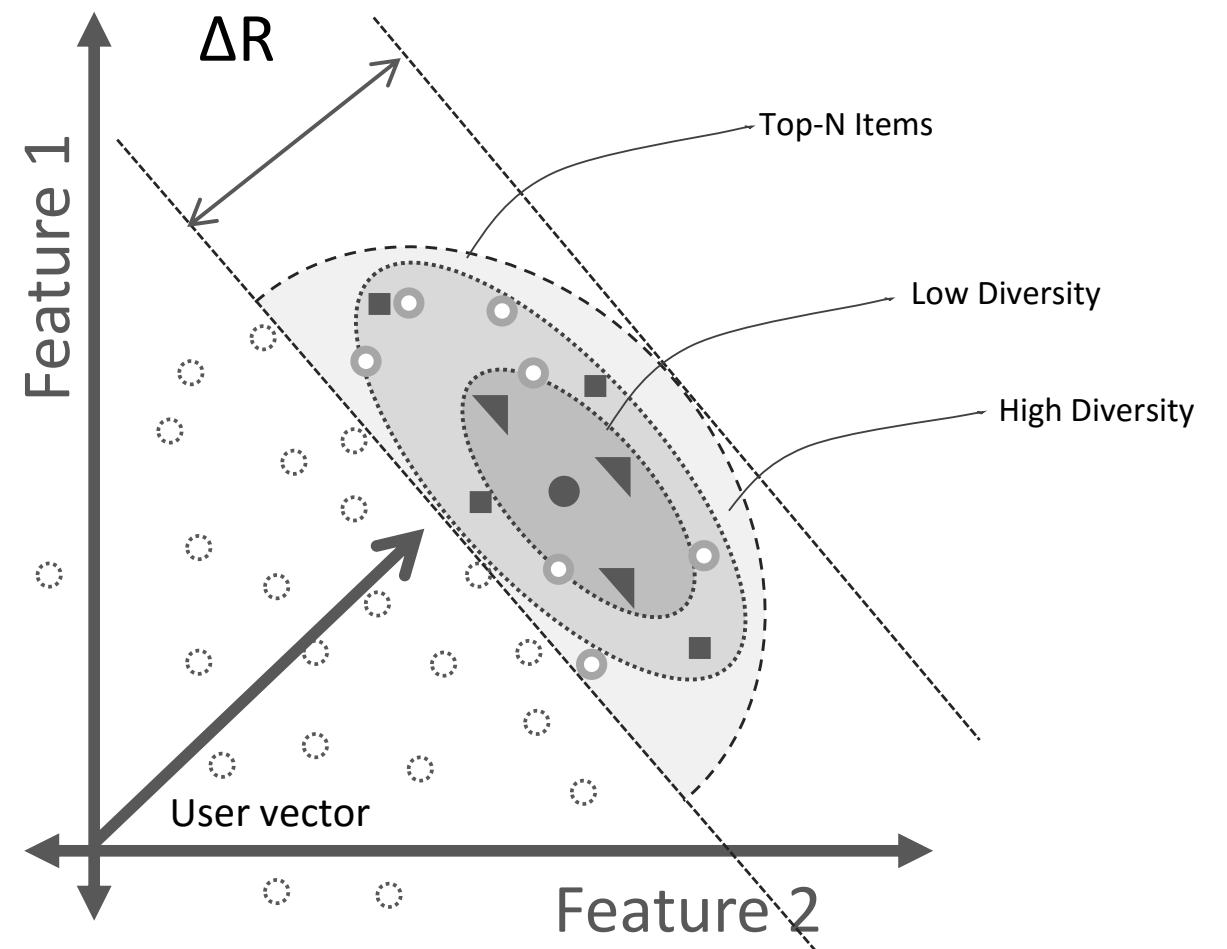


Figure 2. A simplified illustration of the latent factor approach, which characterizes both users and movies using two axes—male versus female and serious versus escapist.





Serendipity

Offline algorithms are tested on items that are **already rated** by the user

To some extent, you are optimizing/testing for “ratability”

How to produce (or evaluate) recommendations that are “non-obvious”?

e.g. Items in the “long tail”

The importance (and “flavor”) of this may depend on the task at hand!



User experience

User satisfaction does not always correlate with high recommender accuracy (McNee 2002, Ziegler 2005)

Why not?

- Importance of other recommendation qualities (diversity, serendipity, etc.)

- Importance of interaction mechanism usability (preference elicitation, transparency, etc.)

- Interpersonal (e.g. cultural) differences

- Situational (e.g. task-related) differences



Contributions

What does this paper contribute to the state-of-the-art?

Mindset: think about algorithms beyond accuracy

Agenda-setting: outline other metrics and objectives for optimize/test algorithms

What else?



Criticisms

What is missing from this paper?

Metrics: how do we measure diversity, serendipity, UX?

Reconciliation: If the metrics disagree, how do we reconcile them?

What else?



User experience

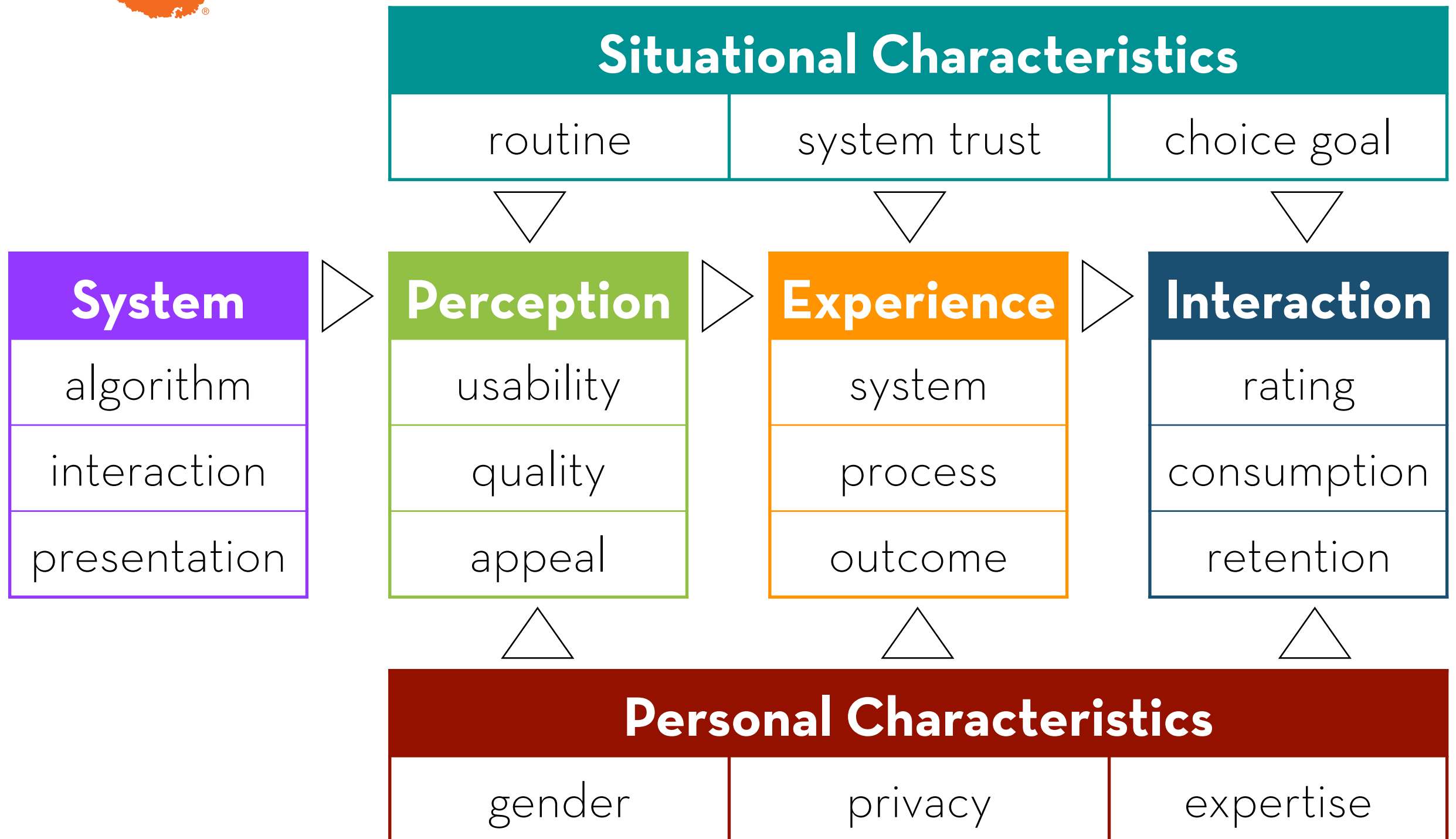
Present a framework to integrate user experience evaluations

Provide measurement scales to measure subjective aspects and experiences

Validate the framework in a series of field trials

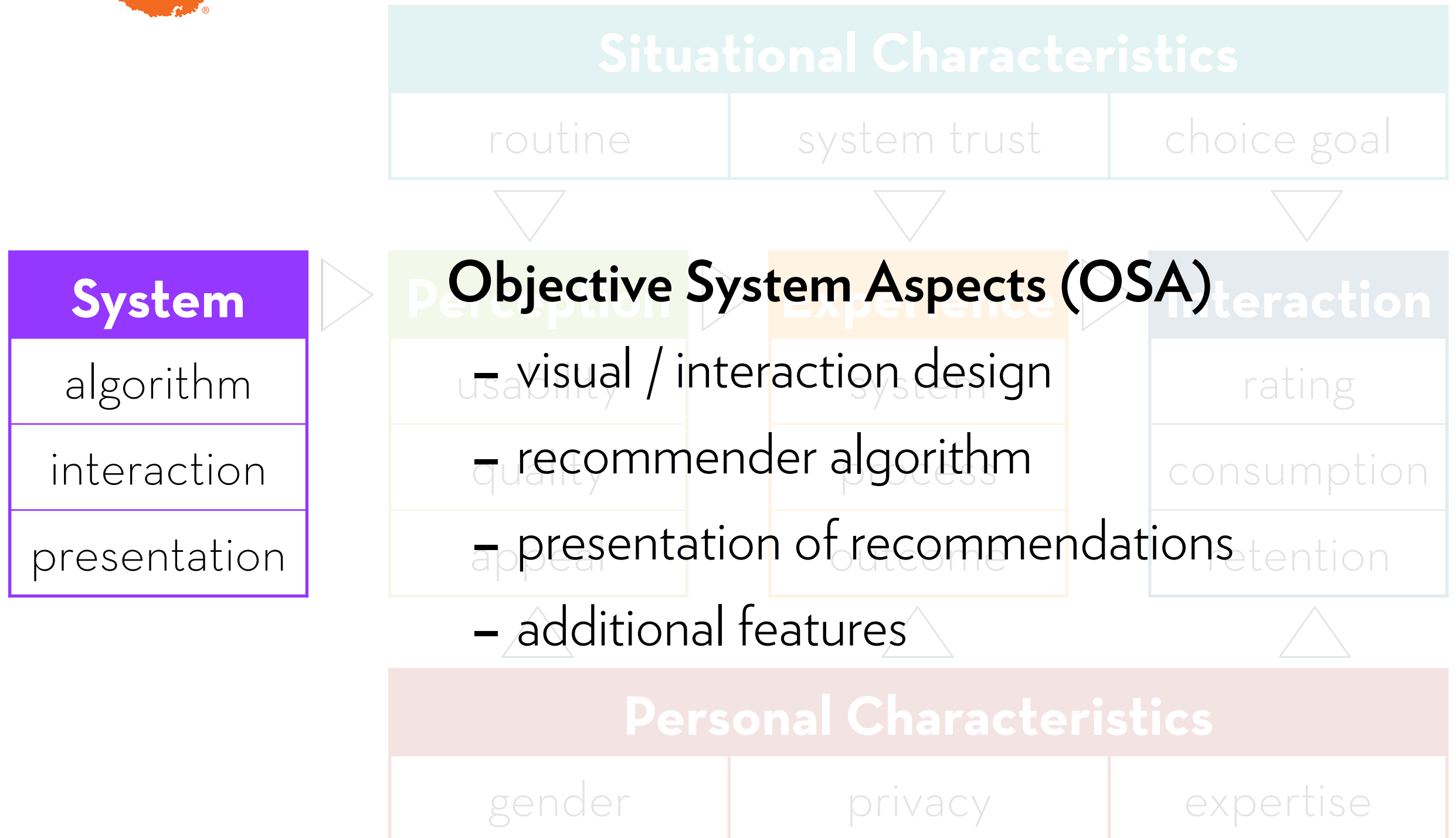


Framework





Framework



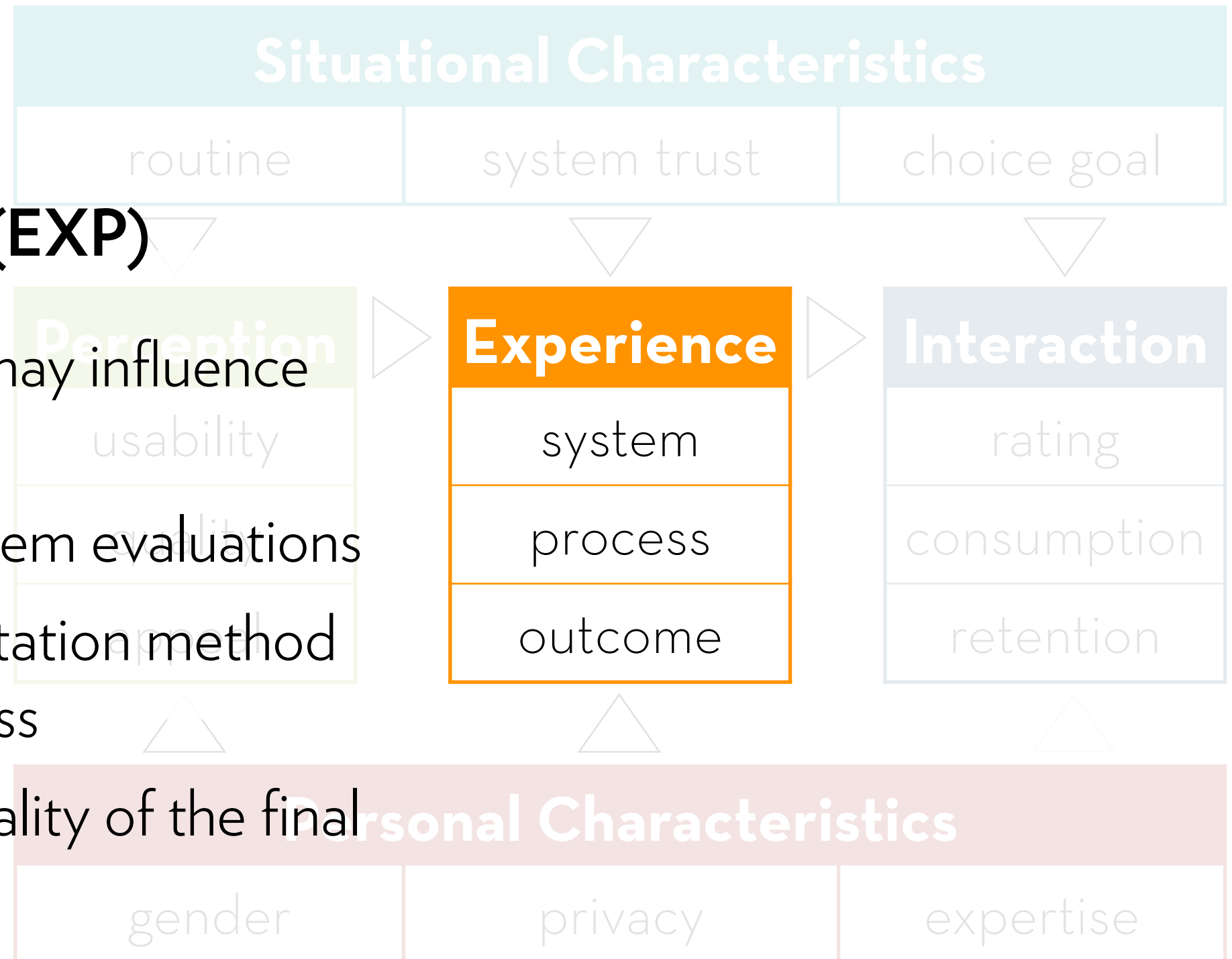


Framework

User Experience (EXP)

Different aspects may influence different things

- interface -> system evaluations
- preference elicitation method
-> choice process
- algorithm -> quality of the final choice





Framework

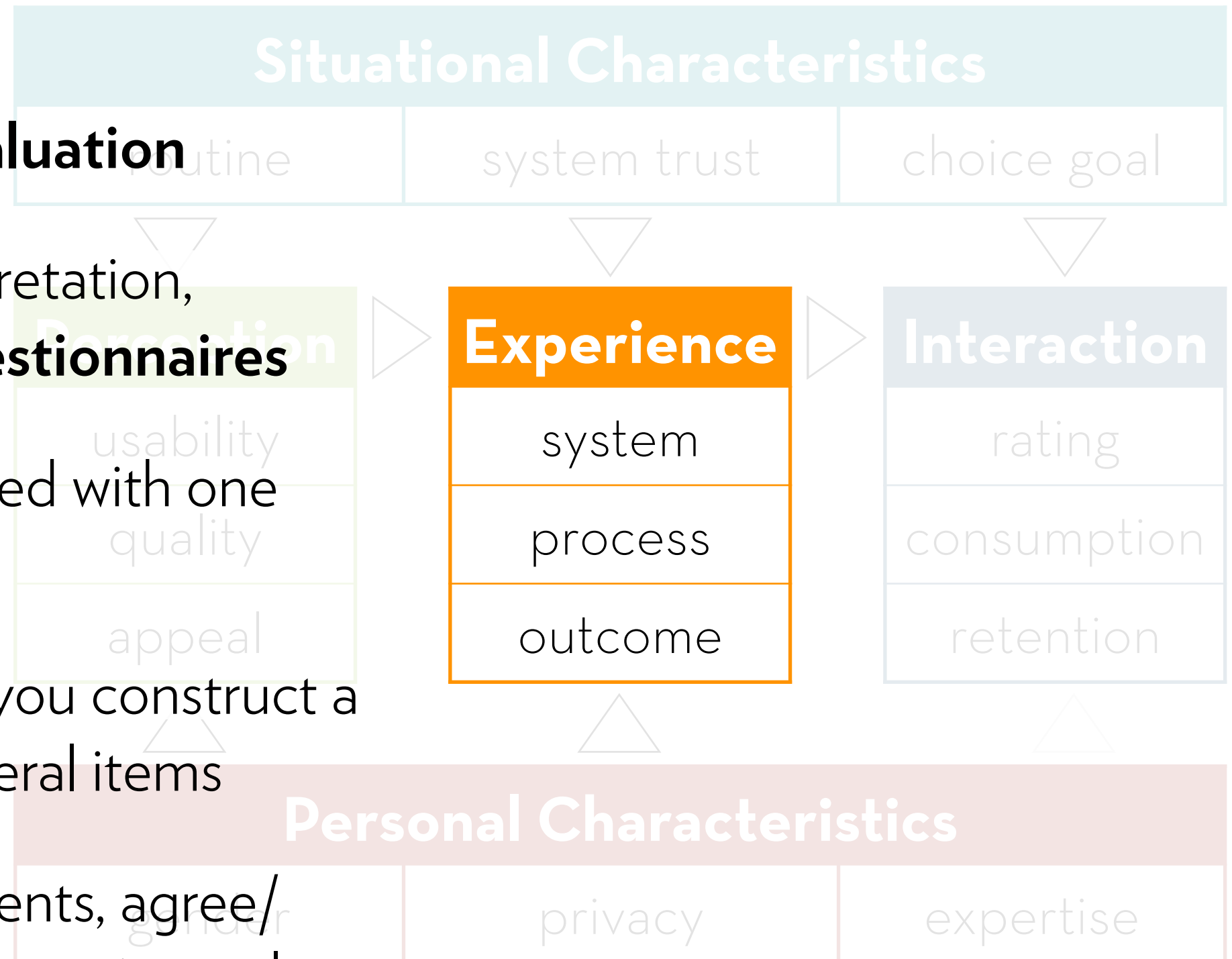
An attitudinal evaluation

self-relevant interpretation,
measured with **questionnaires**

Cannot be measured with one
question

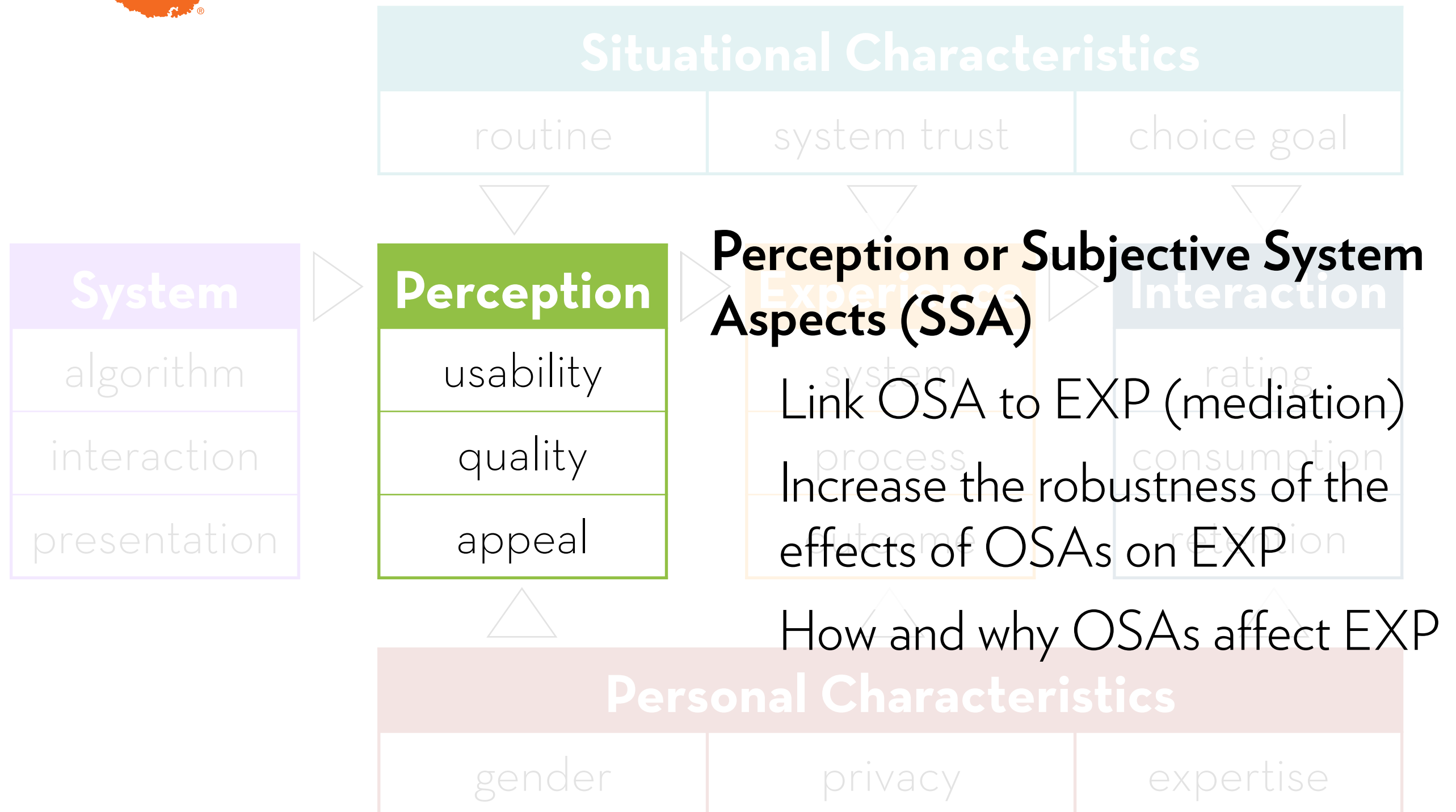
Far more robust if you construct a
scale based on several items

Usually 5-8 statements, agree/
disagree on a 5- or 7-point scale



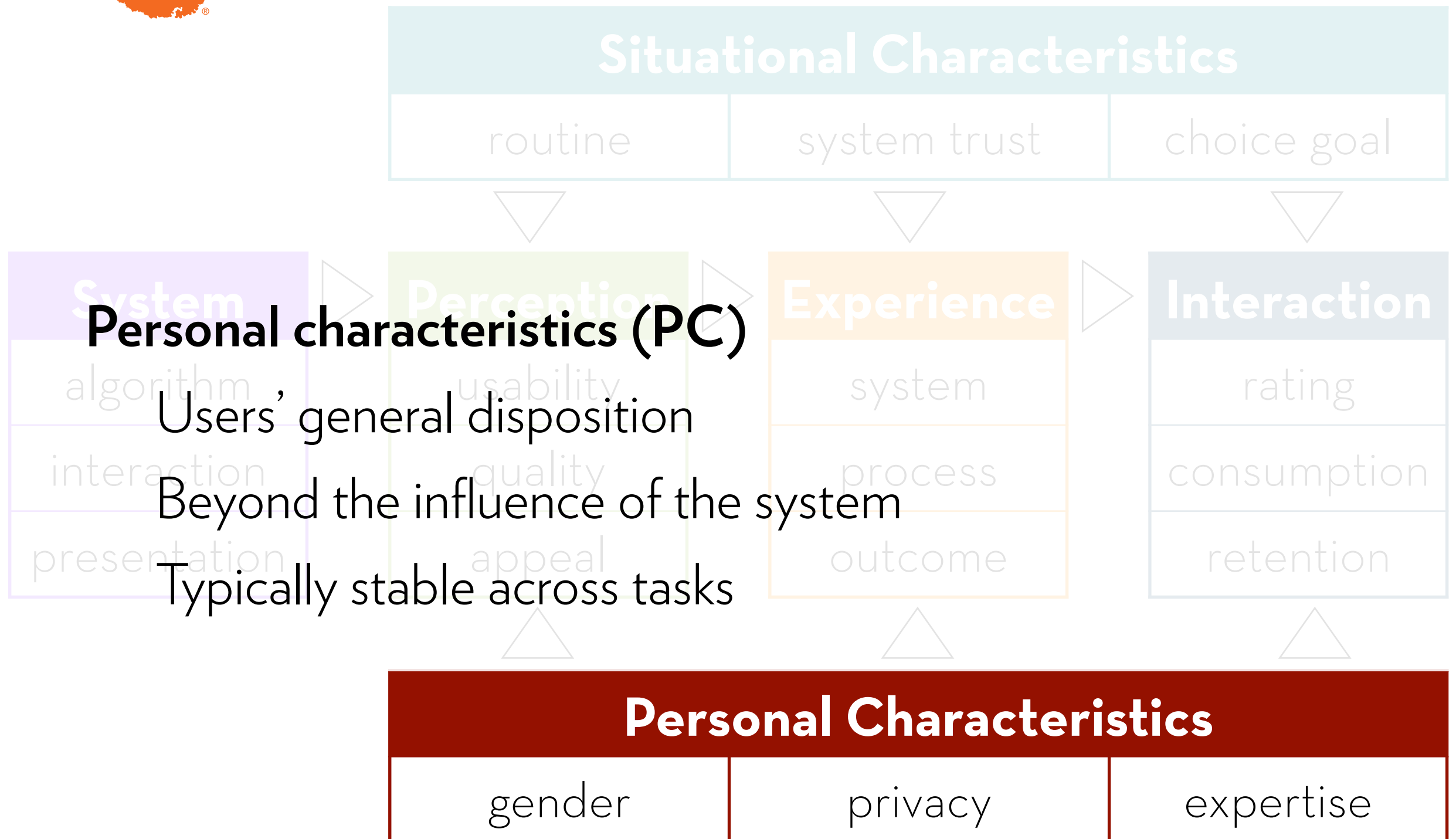


Framework





Framework





Framework

Situational Characteristics

routine

system trust

choice goal

System

Situational Characteristics (SC)

How the task influences the experience

Also beyond the influence of the system

Here used for evaluation, not for augmenting algorithms!

Perception

Experience

Interaction

algorithm

interaction

presentation

usability

quality

appeal

system

process

outcome

rating

consumption

retention

Personal Characteristics

gender

privacy

expertise



Framework

Situational Characteristics

routine

system trust

choice goal

Interaction (INT)

System

Observable behavior

– browsing, viewing, log-ins

Final step of evaluation

But not the goal of the evaluation

Triangulate with EXP

Perception

usability

quality

appeal

Experience

system

process

outcome

Interaction

rating

consumption

retention

Personal Characteristics

gender

privacy

expertise



The full cycle

Link the objective interaction (purple) to objective system aspects (blue)

through a series of subjective constructs

Why? Because these constructs...

...are likely to **attenuate** and qualify the effects

...**explain** why users' behavior is different for different systems

This explanation is the main value of our framework



How to use the framework

What it is not

a **predefined metric** for standardized “performance” tests

What it is

conceptual definition of a generic **chain of effects**

helps researchers to explain **how and why** the user experience of recommender systems comes about

a generic **guideline** for in-depth empirical research

generator for **new hypotheses**



EMIC pre-trial (FT1)

Questions

Do users want to receive personalized recommendations **at all**?

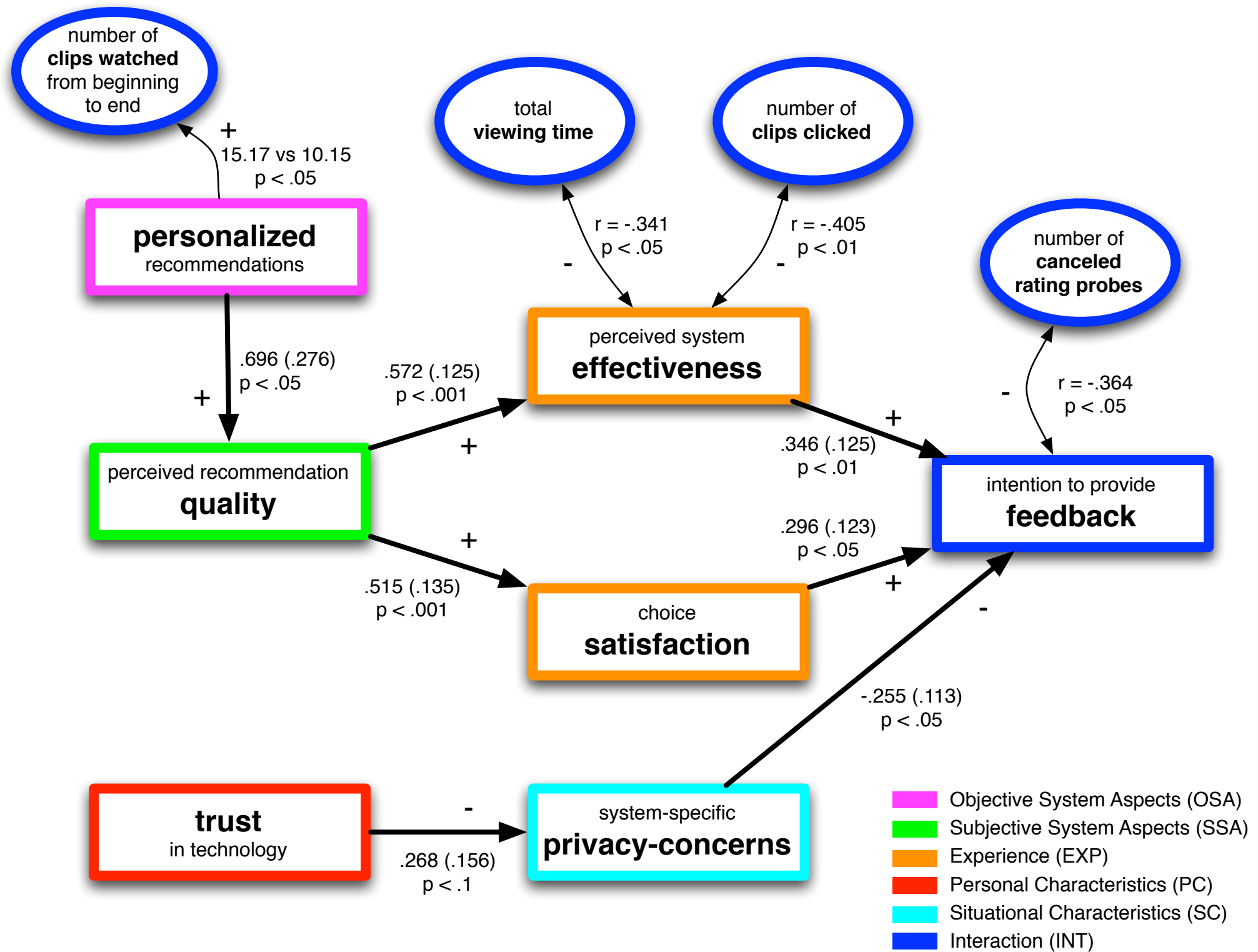
Will users provide adequate preference **feedback**?

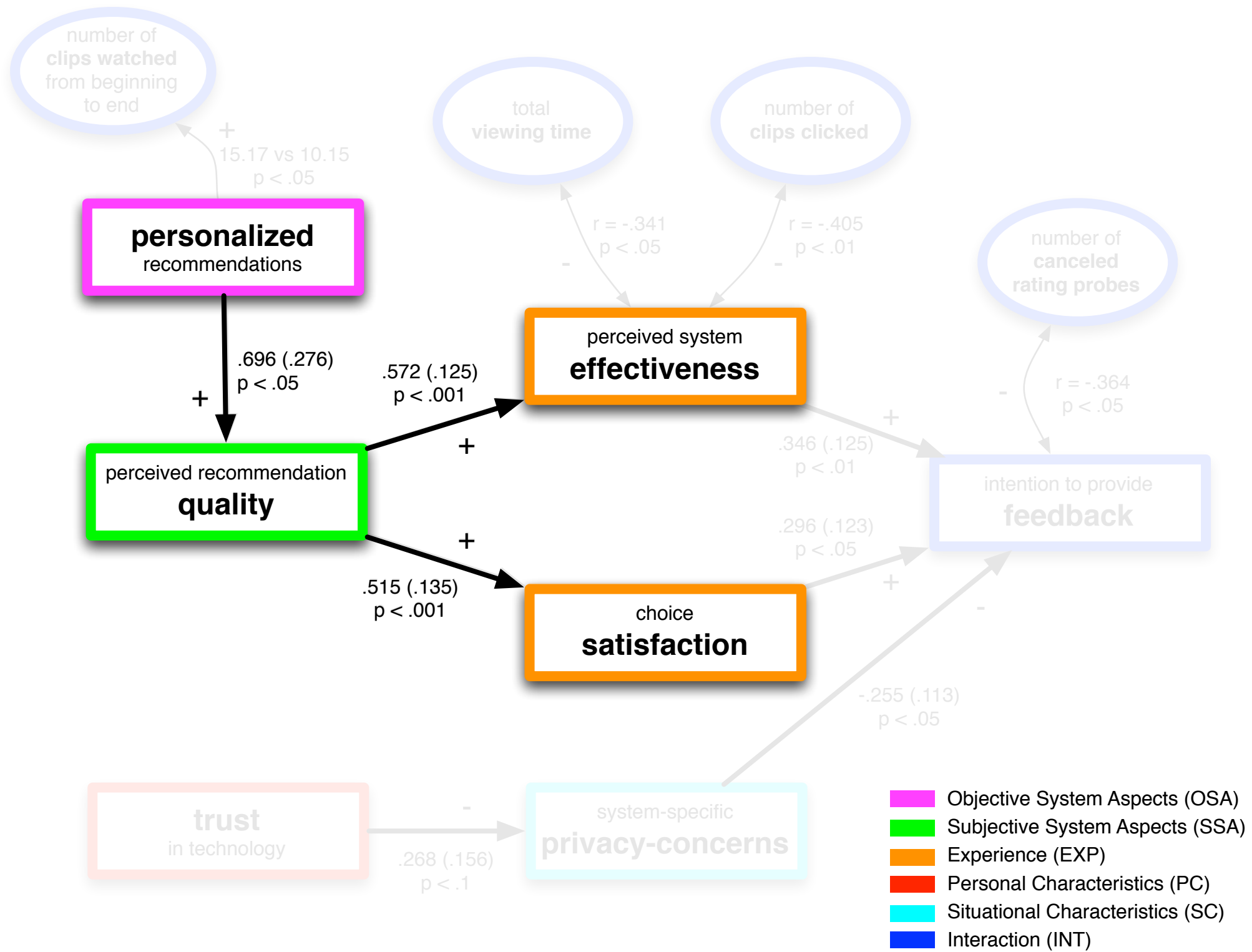
Setup

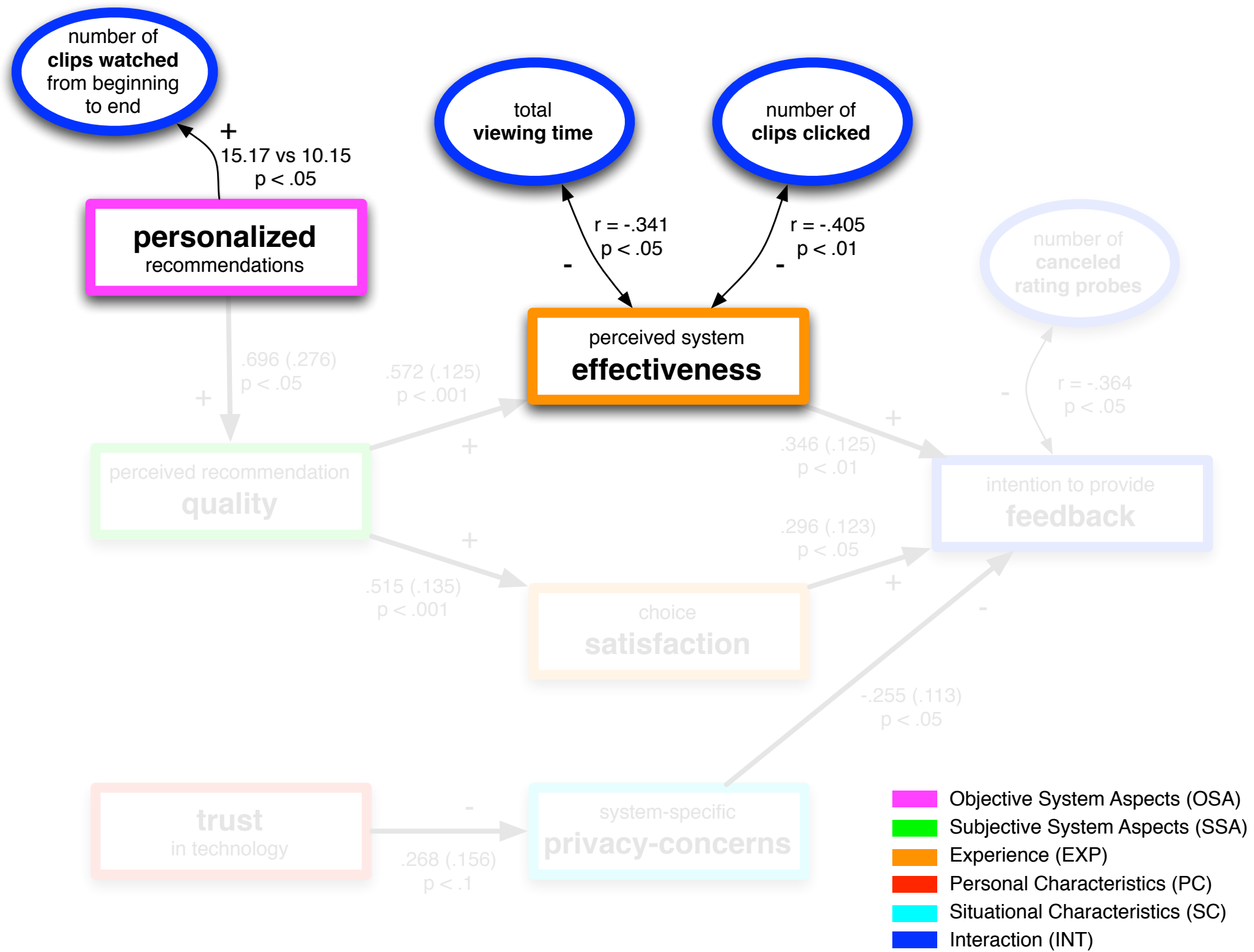
43 participants (between subjects)

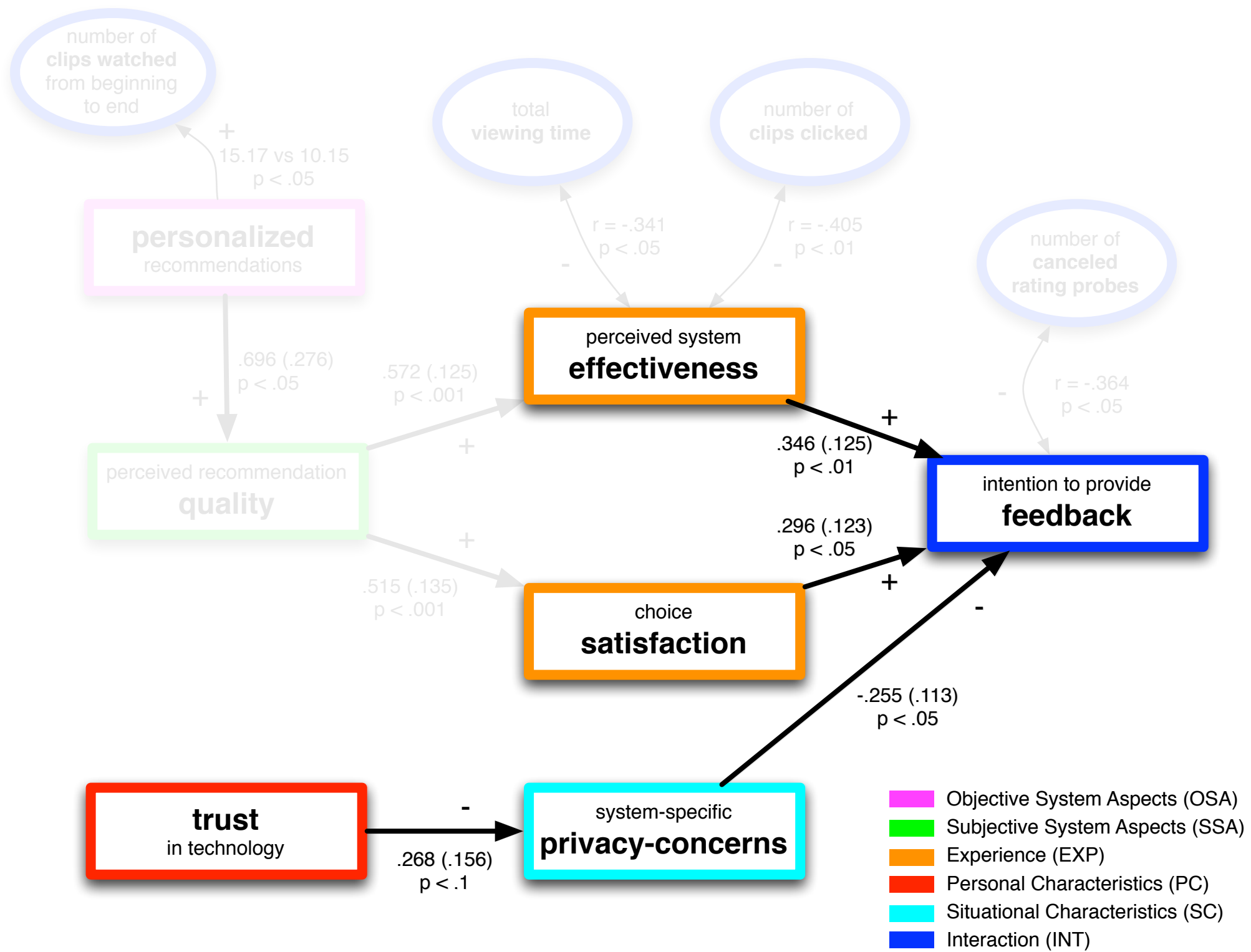
Random recommendations vs. personalized VSM algorithm recommendations (between subjects)

Measure quality perceptions, experience, privacy concerns, and intention to provide feedback; capture clickstream











BBC trial (FT4)

Questions

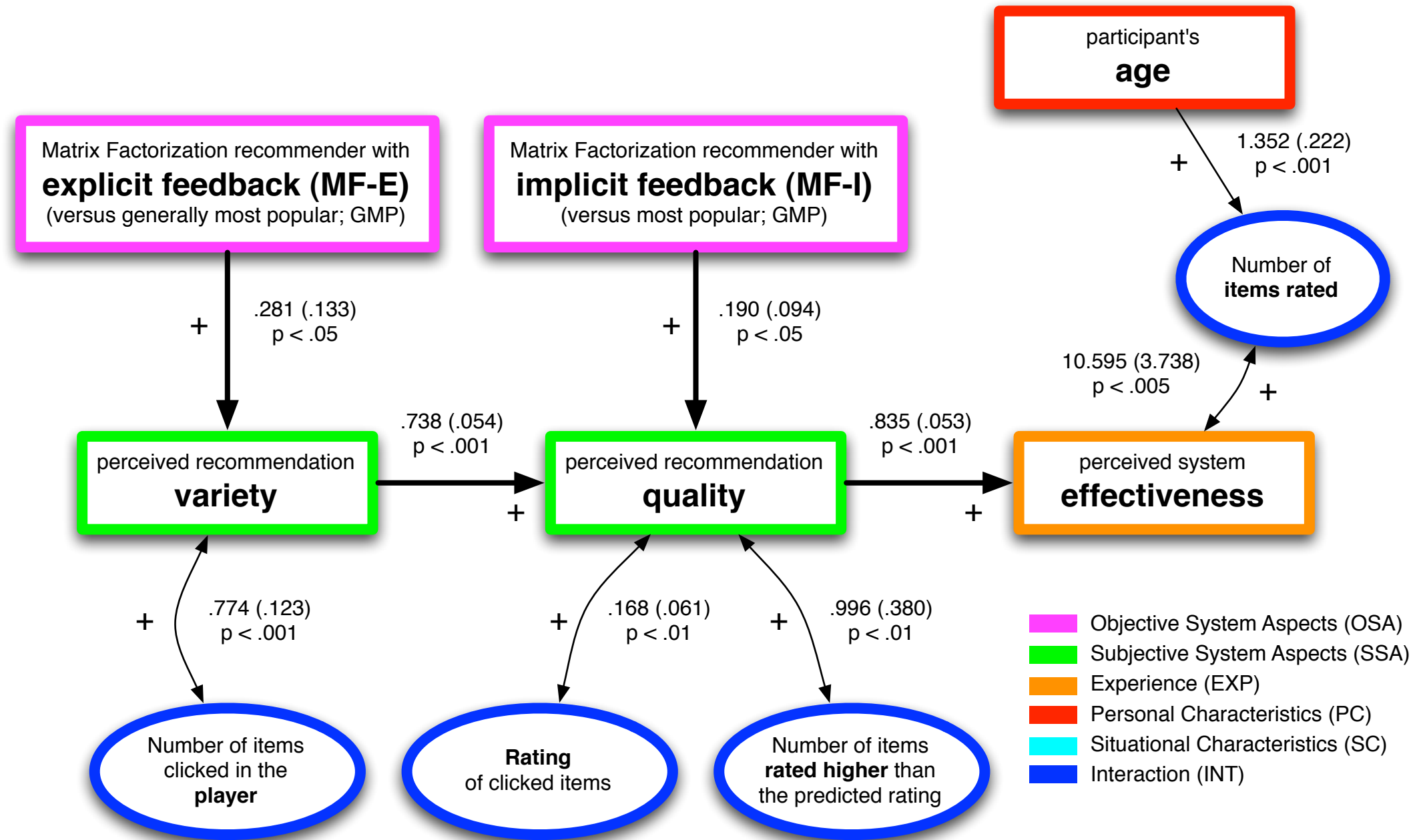
Can we find experience differences between non-personalized recommendations, an MF algo with **explicit feedback**, and an MF algo with **implicit feedback**?

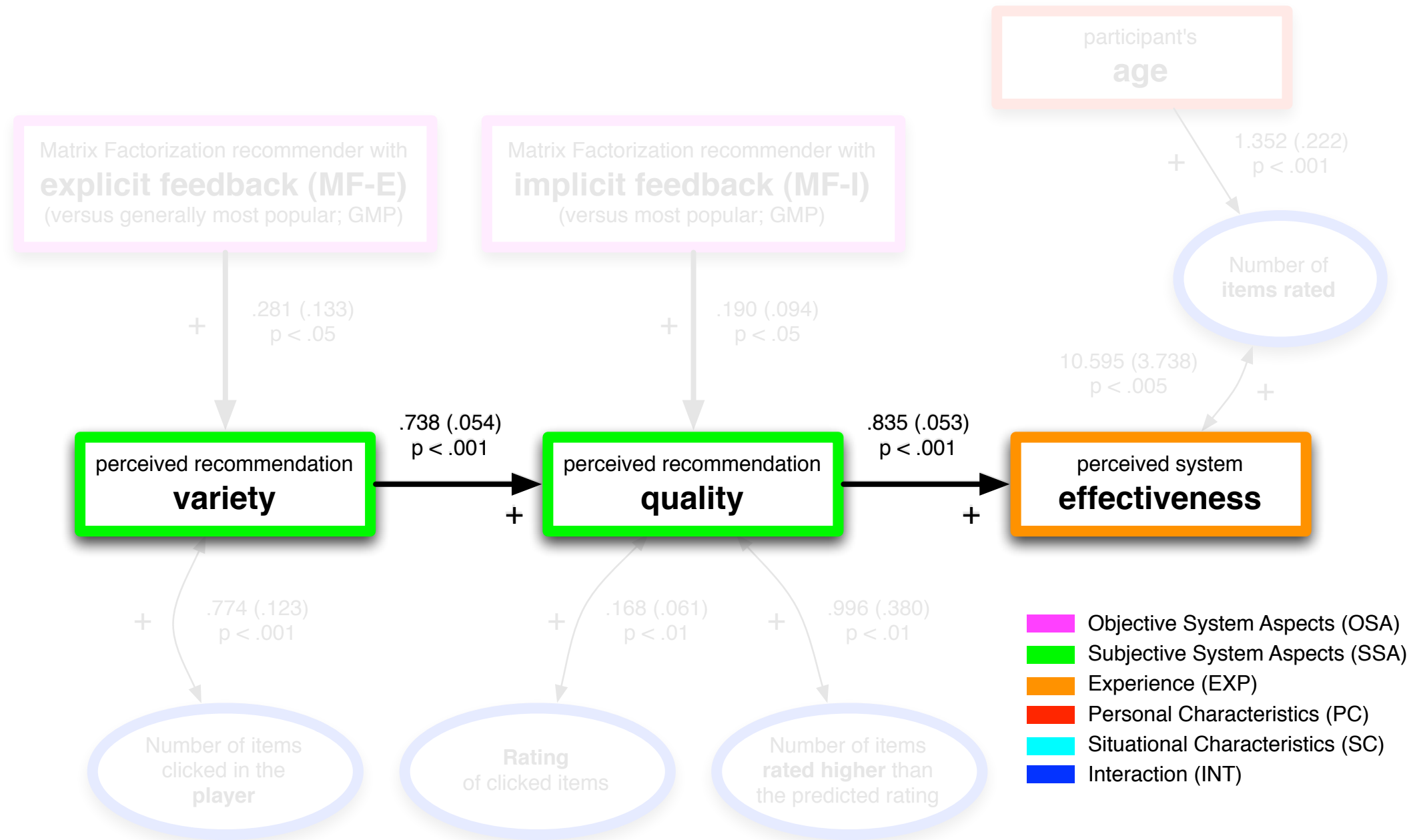
Setup

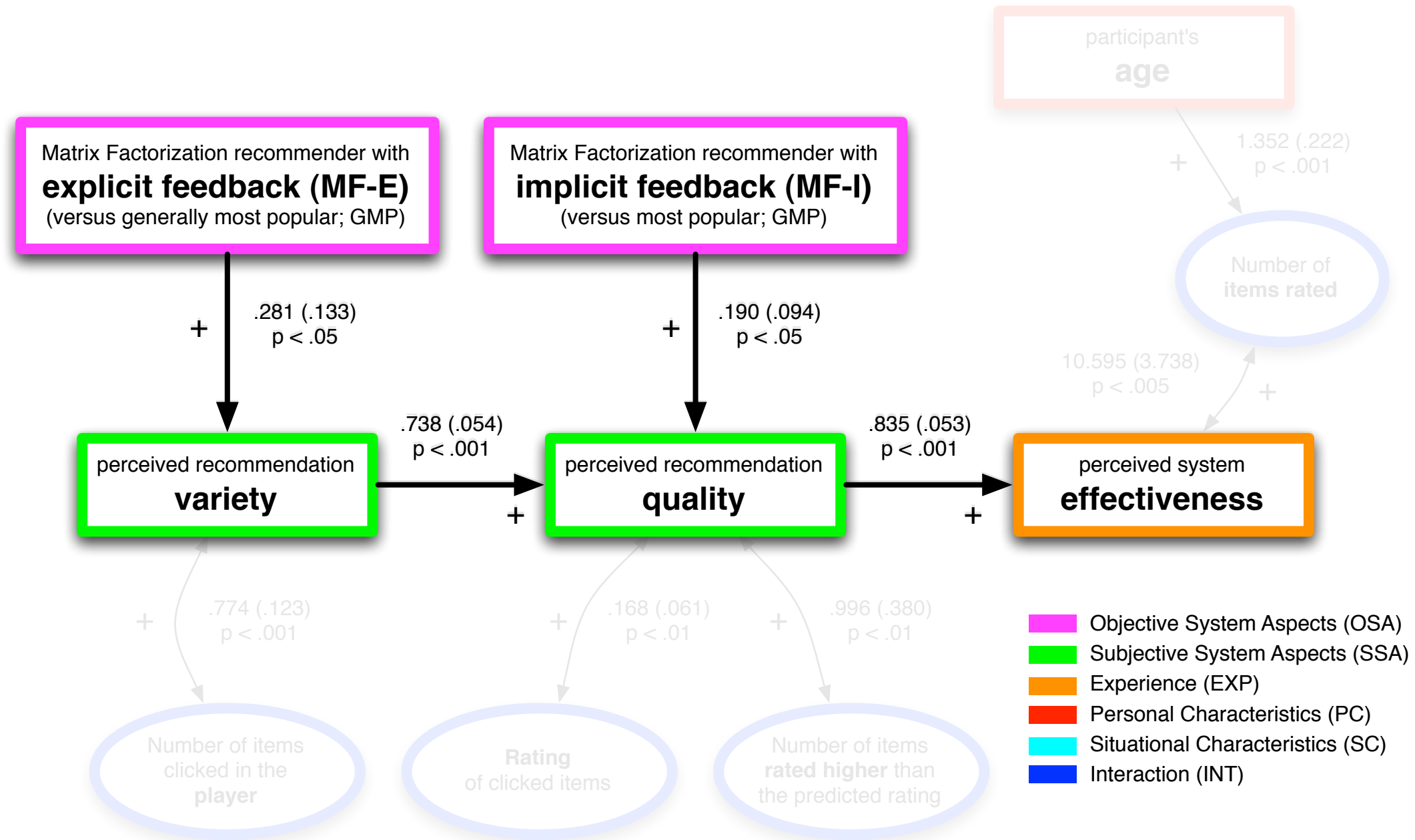
58 participants

Manipulate algorithm/PE method (within subjects)

Measure perceived quality and variety of recommendations, and experience; capture clickstream









Contributions

The **structure** of the framework is validated

With a few exceptions

We developed a number of **measurement scales** to measure SSAs, EXPs, PCs and SCs

More have been added in subsequent work

We covered some **gaps** in existing research

More details are in the paper itself



Questions...