

Foyr Ideate

AI design exploration & selection flow · Cognitive load, trust gaps, decision friction

4 Critical Issues	3 High Issues	2 Medium Issues
----------------------	------------------	--------------------

Core Finding

 **The product generates outputs, not experiences — and users pay the cost**

Foyr Ideate has optimized for AI generation speed while leaving the entire human side of the loop undesigned. There is no scaffolding for evaluation, no feedback mechanism that matches how taste actually works, and no visible reasoning from the AI — so users enter a paralysis loop: too much to look at, no way to react, no path forward that feels safe. The interaction model is a gallery, but users need a collaborator.

Issue Analysis

COGNITIVE LOAD ISSUES

CRITICAL	Simultaneous multi-design display exceeds working memory capacity
ISSUE Multiple AI-generated designs are shown at once with equal visual weight, no	

hierarchy, and no anchor point — forcing users to hold, compare, and evaluate everything simultaneously.

WHY

Miller's Law meets aesthetic complexity: working memory caps at ~4 items, but each design is not a discrete 'item' — it is a dense aesthetic system with dozens of variables. The display treats outputs as thumbnails when users experience

nce
them as
rooms.
The
cognitiv
e load is
not 4x,
it is
expone
ntial.

FIXES

- **S
e
q
u
e
n
t
i
a
l
r
e
v
e
a
l
—
s
h
o
w
1
,
t
h
e
n
l
e
t
u
s
e
r
r
e
q
u
e
s
t
m
o**

- **Anchor + compare : fix on the result, overlay alternative Dimno**
- **Dimno**

n - focused cards on hover

- Limit initial view to 2 options max

CRITICAL No evaluation framework — users don't know what “better” means here

ISSUE

The interface presents competing designs without surfacing comparison dimensions — users have no axis (mood, light, style density, colour temperature) on which to anchor judgment, so they default to paralysis or arbitrary selection.

WHY

The system confuses visual legibility with evaluative legibility. Users can see the designs clearly — but

'seeing'
a room
and
'evaluati
ng' a
design
decision
are
cognitiv
ely
different
tasks.
Without
a
scaffold
ed
evaluati
on
model,
users
apply
none,
freeze,
and
blame
themsel
ves for
not
knowing
what
they
want.

FIXES

- **A
u
t
o
-
g
e
n
e
r
a
t
e
d
c
o
m
p
a
r
i**

s
o
n
t
a
g
s
p
e
r
d
e
s
i
g
n
(
“
w
a
r
m
e
r
p
a
l
e
t
t
e
”
,
“
m
o
r
e
o
p
e
n
p
l
a
n
”
)
• S
i
d
e
-
b

y - s i d e d i f f m o d e w i t h h i g h l i g h t e d d e l t a s

- G u i d e d e v a l u a t i o n p r o

m
p
t
:
“
W
h
i
c
h
f
e
e
l
s
m
o
r
e
l
i
k
e
y
o
u
?”

- D
i
m
e
n
s
i
o
n
s
l
i
d
e
r
s
t
h
a
t
r
e
f
l
e
c

t
c
u
r
r
e
n
t
d
e
s
i
g
n
s
t
a
t
e

TRUST GAPS & UNCERTAINTY POINTS

CRITICAL	AI rationale is invisible — outputs feel arbitrary, not generated
ISSUE Designs appear with no trace of the inputs, logic, or style signals that produced them. Users cannot tell if the AI understood them, guessed randomly, or generated from defaults.	

WHY

The Gulf of Evaluation is total. Norman's principle states that users need feedback to understand whether their actions had the intended effect. Here, the 'action' (entering preferences) is disconnected from the 'effect' (designs generated) with zero visible trace — so users can't evaluate the AI's understanding, only its output. Trust is impossible to build

without
a
reasoni
ng trail.

FIXES

- “
**W
h
y
t
h
i
s
d
e
s
i
g
n
?**
”
c
a
r
d
u
n
d
e
r
e
a
c
h
r
e
s
u
l
t
- I
n
p
u
t
e
c
h
o
:
“
**G
e**

generated from: warm tones · open layout · Nordic style”

- Highlight

Which user inputs influenced which design elements?
Confidence

s
i
g
n
a
l
:
“
S
t
r
o
n
g
m
a
t
c
h
”
v
s
“
C
r
e
a
t
i
v
e
i
n
t
e
r
p
r
e
t
a
t
i
o
n
”

HIGH

Input-to-output disconnect makes users feel like spectators, not authors

ISSUE
Users
feel like

they are
guessin
g when
they
interact
with the
AI —
there is
no
visible
mappin
g
betwee
n what
they
specifie
d and
what
was
produce
d,
destroyi
ng the
sense
of
creative
agency.

WHY

The
Gulf of
Executi
on is
wide
and
unackn
owledge
d. Users
need to
feel that
their
intent
maps to
system
behavio
r. When
the
system
acts like
a black
box,
users
experie
nce it as
a slot

machine —
and slot
machine
interactions are
engaging for
gambling, not
for
high-investment
design
decisions where
authorship
matters.

FIXES

- Live
previews
w/resp
onses
as
users
adjust

inputs
• Intent confirmation (“Here, what I, m

designing for you ...

-) Input - tagged design elements (tap elements

t
→
s
e
e
w
h
i
c
h
i
n
p
u
t
d
r
o
v
e
i
t
)

DECISION-MAKING FRICTION

CRITICAL	Feedback mechanism is absent — refinement requires language users don't have
ISSUE Users struggle to give feedback to refine designs. The only apparent path is re-entering text prompts — but design preference is pre-verbal and users can't	

articulate what's wrong with a room, only that something is.

WHY

The feedback model is built for designers, deployed to homeowners. Expert critique requires vocabulary that most users have never developed. Asking them to type feedback is the equivalent of asking someone to describe the taste they want before they've tasted anything. The refinement loop

needs
input
modaliti
es that
match
how
non-exp
ert taste
actually
works:
visual,
compar
ative,
gestural
.

FIXES

- **B
i
p
o
l
a
r
s
l
i
d
e
r
s
:
w
a
r
m
e
r
↔
c
o
o
l
e
r
,
m
i
n
i
m
a
l
↔
l**

a
y
e
r
e
d
,
b
r
i
g
h
t
↔
m
o
o
d
y
T
a
p
-
t
o
-
a
n
n
o
t
a
t
e
:
m
a
r
k
w
h
a
t
y
o
u
l
i
k
e
/
d
i
s

I
l
i
k
e
o
n
t
h
e
i
m
a
g
e
s

- S
w
i
p
e
c
a
r
d
r
a
t
i
n
g
:
k
e
e
p
/
l
o
s
e
/
m
o
r
e
l
i
k
e
t
h
i
s

- A
l

- suggested refinements: "Want it to consider? Try this variation →"

CRITICAL	Selection feels irreversible — fear of commitment triggers avoidance
ISSUE Users hesitate to select a design because they fear making the wrong choice. There are no visible undo affordances, no save-for-later, no 'shortlist' model — selection signals finality. WHY Loss Aversion in a domain with no safety net. Interior design carries high real-world stakes in the user's mind, even at the explorat	

ion
phase.
When
an
interfac
e
provide
s no
reversib
ility
signal, it
amplifie
s that
perceiv
ed risk.
The UX
should
commu
nicate
“this is a
canvas,
not a
contract
” — but
instead
it reads
like a
checkou
t
screen.

FIXES

- **E
x
p
l
i
c
i
t
r
e
v
e
r
s
i
b
i
l
i
t
y
l
a**

b e l : “ Y o u c a n a l w a y s c h a n g e o r r e g e n e r a t e ”

- **S h o r t l i s t / b o a r d m o d**

el : save multiple people options without committing
• Versioning history :

a
u
t
o
-
s
a
v
e
e
v
e
r
y
d
e
s
i
g
n
s
t
a
t
e
C
T
A
c
c
o
p
y
s
h
i
f
t
:
“
E
x
p
l
o
r
e
t
h
i
s
d
i
r
e

c
t
i
o
n
→
”
n
o
t
“
S
e
l
e
c
t
”

USABILITY ISSUES

HIGH	No progress model — users don't know where they are in the design journey
ISSUE The flow lacks a visible design journey model — users don't know if they're at the beginning of exploration, mid-refinement, or approaching a final decision. Session s feel	

undirected.

WHY

Without a map, every action feels equally meaningless.

Visibility of system status (Nielsen heuristic #1)

applies not just to loading states but to process position.

Users need to know: am I narrowing down or still exploring? That orientation changes how they make decisions.

FIXES

- **J
o
u
r
n
e
y
s**

target indicator: Explore → Refine → Commit

- Session bread crumb: s

how how designs have evolved

- Milestone prompts: “You’ve shortl

i
s
t
e
d
3
—
r
e
a
d
y
t
o
c
o
m
p
a
r
e
?
”

MEDIUM	Design results lack hierarchy — everything demands equal attention
ISSUE All generat ed designs are present ed with identical visual weight — same size, same promine nce, no AI-rank ed ‘best match’ signal. Users have no starting point and no natural	

focus of
attention.

WHY

Visual democracy in a results set is a UX abdication. When every result looks equally important, the UI delegates the entire evaluation burden to the user. A ranked or signalled hierarchy — even a 'closest to your brief' badge — gives users an anchor and dramatically reduces time-to-first-interaction.

FIXES

- “
B

estimatich" badge on highest - confidence result

- Primary + second

array layout : 1 heroresult + 2 alternatives
• Sort options : by match

%
,
b
y
s
t
y
l
e
,
b
y
m
o
d

HIGH	No progress model — missing journey stage indicator
ISSUE Users don't know whether they are in an early exploration phase, a refinement phase, or approaching commitment. Without stage clarity, every decision feels equally weighty. WHY Lack of wayfinding collapse	

s the mental model of the task. When users can't locate themselves in a process, they treat every action as high-stakes, amplifying hesitation across the entire flow.

FIXES

- **E
x
p
l
i
c
i
t
E
x
p
l
o
r
e
→
R
e
f
i
n
e
→
C
o
m**

mit progress indicator Stage - contextual prompts (“Still exploring

o r i n g ? S a v e a f e w f a v o u r i t e s f i r s t)

- A u t o - a d v a n c e s t a g e w h e n s h

o
r
t
l
i
s
t
t
h
r
e
s
h
o
l
d
i
s
m
e
t

A/B Experiments

A/B TEST 1 Rationale card vs. no rationale on design outputs		
CONTROL Designs shown with no explanation — current state	VARIANT “Why this design?” card under each result showing input echoes and style tags	METRICS Selection rate · Time-to-first-click · Regeneration rate
A/B TEST 2 Sequential reveal (1 at a time) vs. grid display (all at once)		
CONTROL Full grid of generated designs shown simultaneously	VARIANT One design shown at a time with “Show another” and “Keep this” actions	METRICS Selection confidence score (post-task) · Session completion · Drop-off rate

A/B TEST 3 Structured micro-feedback (sliders) vs. freeform text refinement		
CONTROL Text prompt input to refine design — current state	VARIANT Bipolar sliders (warmer/cooler, minimal/layered) + “More like this” / “Less like this” tap	METRICS Refinement engagement rate · Iterations per session · Final selection rate

A/B TEST 4 “Explore this direction” CTA copy vs. “Select” CTA		
CONTROL Primary CTA labeled “Select” or “Choose this design”	VARIANT CTA labeled “Explore this direction →” with secondary “Save to shortlist” action	METRICS CTA click-through rate · Bounce from CTA · Downstream engagement depth

What to Fix First

1	Fix Now Surface AI rationale on every design output. This is a one-sprint intervention with outsized return: it directly addresses trust collapse, reduces the evaluation burden, and turns an opaque output into a legible design decision. Without it, every other improvement is built on a broken foundation.
2	Next Sprint Replace freeform text refinement with structured micro-feedback (sliders + tap interactions). The current feedback model is inaccessible for the vast majority of users. Non-verbal feedback mechanisms will unlock the refinement loop entirely and materially increase session depth, return visits, and selection rate.
3	Strategic Redesign the output presentation model from gallery to guided collaboration. Introduce sequential reveal, a shortlist board, reversibility signals, and a journey stage indicator. This is a systemic reframe of the product interaction model — from “AI generates, human picks” to “AI and user co-narrow toward a decision.”