



Assignment1 - Ph3 - Brainstorming & Wireframes

STEPS

Step 1: Brainstorm as many quality solutions as possible (Diverge) - *Generate wide* - 30m

Step 2: Prioritise Smartly (Converge) - *Focus narrow* - 15m

Step 3: Wireframe Your Top 2-3 Ideas, showing user flows - *Make it visual* - 1h 30m

Step 4: Define Success Metrics - *Plan measurement* - 15m

Step 5: Identify Pitfalls & Mitigations - *Anticipate failure* - 15m

1. Brainstorm Solutions

Problem Statement:

"Users researching new products on Amazon lack the context to evaluate which product features are critical for their specific needs versus which are irrelevant or marketing-driven."

Root Cause: Amazon shows specs but doesn't explain relevance or priority hierarchy for user's needs.

▼ Problem Statement & Details

PROBLEM STATEMENT (Feature Guidance Gap - P0)

Full Definition:

"Users researching new products on Amazon lack the context to evaluate which product features are critical for their specific needs versus which are irrelevant or marketing-driven. This forces users to:

- **Spend 2-4+ hours researching** (vs. 30 min for guided buyers)
- **Create external tools** (Google Sheets, ChatGPT) to synthesize information
- **Abandon research** or make low-confidence purchases
- **Buy suboptimal products** (picking mid-range "safe" options instead of best-fit products)
- **Experience higher returns** (wrong product choice = dissatisfaction)

Root Cause: Amazon shows specs but doesn't explain relevance or priority hierarchy for user's needs.

Affected Users:

- Shivam (28, PMM): "Comparison of 2-3 products is confusing"
- Bharanee (34, Developer): "Amazon description too wordy, wants key points"

- 30 respondents (67%) cited need for feature guidance
- 47% confused by "too many options" or "which features matter"

Impact on Conversion Goal:

- Primary lever for 15% conversion target
 - Reduces decision time 50% (1-2 hrs → 30 min)
 - Increases conversion confidence → Reduces returns
 - Moves buyers from mid-range to optimal products (AOV uplift 5-8%)"**
-

EVIDENCE SUPPORTING P0 CHOICE

From Survey Data:

-  30% explicitly request "help me understand features"
-  47% confused by "too many options" or "which features matter"
-  22% confused by "technical specifications"
-  Research time correlates with feature confusion (2-4+ hrs when unsure about specs)

From Interviews:

-  **Shivam:** "Comparison of 2-3 products is confusing—trying to figure out which one to buy"
-  **Shivam:** "When I come across negative reviews, research takes more time... trying to understand if it matters to me"
-  **Bharanee:** "Amazon description is like a paragraph, while I want to just skim through key points"
-  **Bharanee:** "What I need—Price, Specs, Time to deliver, Ratings" (in order)
-  **Debdatta:** "Too many variants, getting them compared takes longer"
-  **Ila:** "Selection of the product is confusing because of various brands & variants"

From Personas:

-  **Shivam's blocker:** "Which of these 3 is BEST for me?" (comparison confusion due to feature importance unknown)
 -  **Bharanee's blocker:** Feature clarity + plain-English specs
-

Problem Statement for Engineering:

"Users don't understand which product features are critical for their needs, causing 2-4+ hour research times, decision paralysis, and suboptimal purchases. Solve by guiding users through feature priority based on their use case."

Success Metric:

"Reduce average research time for new products from 1-2 hours to 30 minutes, improving conversion rate by 7-10%."

Possible Solutions:

- Have a "Know more" icon that opens a pop up to read about the product
 - What features matter most in order of priority
 - What need is met for a particular feature
- Have an AI icon against a product that opens an interactive window (like Rufus) that simply asks additional questions to the user - to understand specific user needs - in a clean, clear and fixed format and accordingly educates the user on what features to look out for. User can minimize or reopen this AI window any number of times, without losing the context & info, unless the app is killed. Option to save this interaction on local system via download.
- Give a color coding on importance of each spec/ feature shown in Product Description - say, something like light green to dark green.

#	Solution	Technique Used	Core Idea
1	Use-Case-Based Decision Tree	How Might We + Journey Mapping	Ask 3-4 quick questions → show only relevant features for their scenario
2	Amazon Buying Guides	Analogy Thinking (Wirecutter model)	Expert-authored guides explaining feature priorities + recommend 3-5 products by persona
3	AI-Scored Feature Importance	Root-Cause Analysis	Scan reviews → score which features matter most → display importance rankings
4	Feature Comparison Matrix	SCAMPER + Journey Mapping	Interactive side-by-side comparison with color-coded tradeoffs + tooltips
5	Three-Spec Summary Card	Constraint Reversal	Show only 3 critical specs in huge text; hide rest behind "Advanced Specs"
6	Category Education Videos	How Might We + Analogy	2-3 min videos explaining "What specs matter and why" (like YouTube explainers)
7	Plain-English Spec Translator	Journey Mapping + Root-Cause	Auto-translate jargon to plain language + real-world context examples
8	Smart Feature-Based Recommendations	SCAMPER + Constraint Reversal	"Products that match your needs" instead of "customers also viewed"
9	Use-Case Presets	How Might We + Analogy	Pre-built profiles (Allergy Sufferer, Budget Buyer, Professional) → auto-suggest products
10	Interactive Decision Flowchart	Root-Cause + Journey Mapping	Progressive questions that narrow options + personalized "Spec Card" with top 3 matches

2. Prioritise Smartly

#1 USE-CASE-BASED DECISION TREE

RICE Score: 13.50

Metrics:

- Reach: 3/3 (affects users confused about features)
- Impact: 3/3 (directly solves core blocker)
- Confidence: 3/3 (47% explicitly request feature guidance)
- Effort: 2/3 (medium - questionnaire logic + tagging)

Why This Wins:

- Directly addresses Shivam's pain: "Comparison of 2-3 products is confusing"
- Directly addresses Bharanee's pain: "Which features matter for my needs?"
- 47% of users explicitly requested feature guidance
- Creates foundation for other solutions (recommendations, presets depend on this)
- Clear ROI: Research time 1-2 hours → 30 min

#2: PLAIN-ENGLISH SPEC TRANSLATOR**RICE Score: 13.50** (Tied with #1)**Metrics:**

- Reach: 3/3 (affects all product pages)
- Impact: 3/3 (removes jargon barrier)
- Confidence: 3/3 (Interview evidence: Bharanee "too wordy, wants key points")
- Effort: 2/3 (medium - content writing + linking)

Why This Wins:

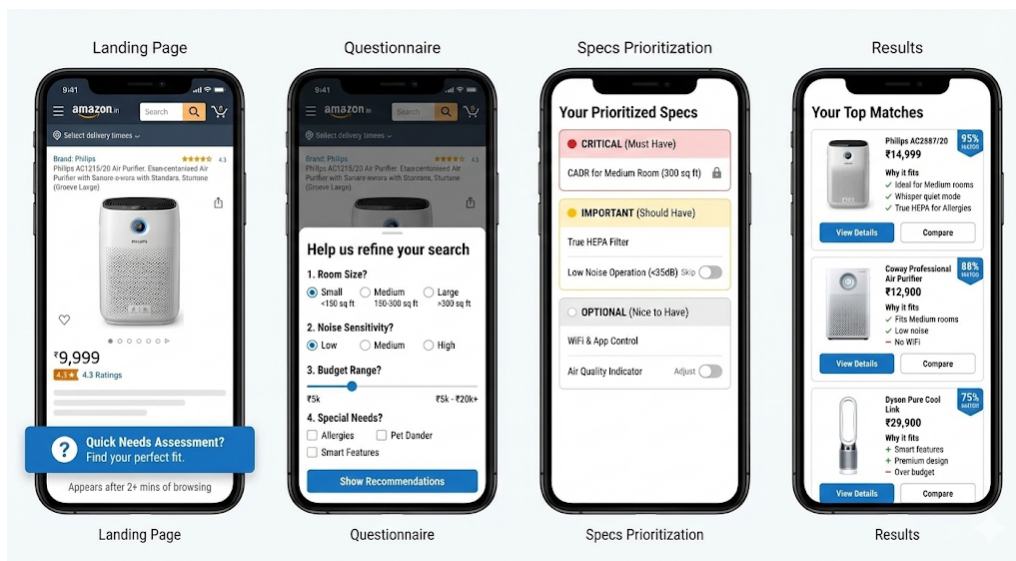
- ✓ Directly addresses Bharanee's pain: "Amazon description too wordy"
- ✓ Affects ALL users (universal benefit)
- ✓ Complements Decision Tree perfectly:
 - Decision Tree: "Which specs matter?"
 - Translator: "What do they mean?"
- ✓ Lower execution risk (simpler than questionnaire logic)
- ✓ Can run in parallel with Solution 1 (different teams)

WHY THIS PAIR (Not Alternatives)**Solution 1 + Solution 7 = Complete Feature Education****RICE SCORES RANKING**

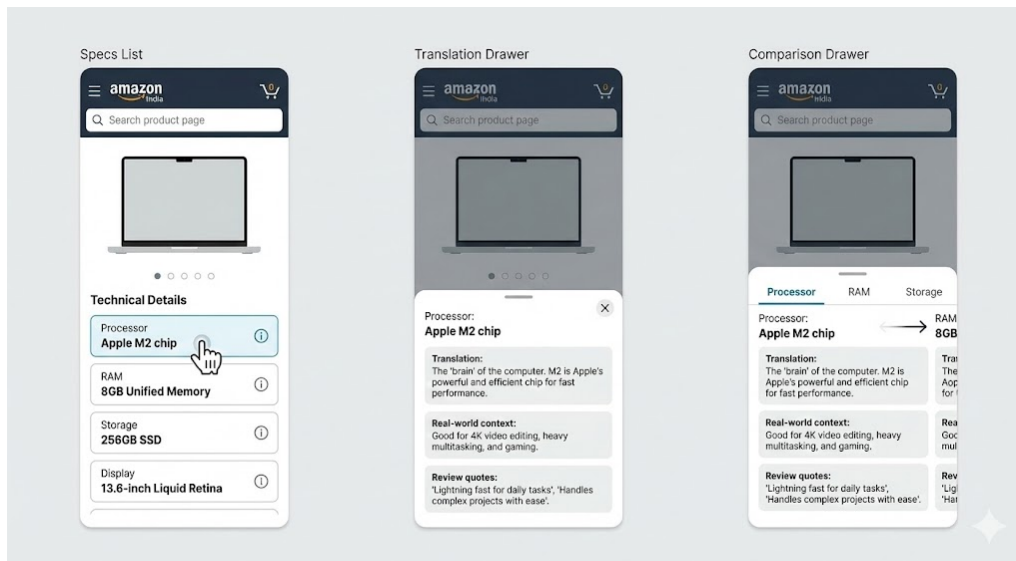
Rank	Solution	RICE Score	Why
1	Solution 5: Three-Spec Summary Card	18.00	Quick win (low effort, high reach)

Rank	Solution	RICE Score	Why
2	Solution 1: Use-Case-Based Decision Tree	13.50	★ TOP PICK (highest impact, medium effort)
3	Solution 7: Plain-English Spec Translator	13.50	★ TOP PICK (tied, complements Solution 1)
4	Solution 4: Feature Comparison Matrix	6.00	Good but secondary
5	Solutions 2, 9, 10 (Guides, Presets, Flowchart)	4.00	High effort, lower reach
6+	Solutions 3, 6, 8 (AI, Videos, Recs)	1.33-2.00	Low confidence or limited reach

3. WIREFRAMES



Solution 1 - Use Case based Decision Tree - Guiding users on Product Research of a new product



Solution 2 - The Specs Translator (to plain English) - helping users understand and compare product variants

4. SUCCESS METRICS

Metric	Baseline	Target	How to Measure
Research Time	1-2 hours (complex products)	30 minutes	Session analytics (time from product discovery to purchase)
Conversion Rate	~22% (estimated baseline)	32%+ (+10-12% lift)	Purchase conversion tracking
Adoption	N/A	60%+ of users try Decision Tree	Feature usage tracking
User Confidence	47% confused about features	<20% confused	Post-research survey: "How confident are you in your choice?"
Spec Understanding	Baseline TBD	+40% better understanding	Post-research survey: "Did you understand what each spec means?"

Last two metrics require surveys after implementation to note the difference in average ratings.

5. RISK & MITIGATION FRAMEWORK

SOLUTION 1: USE-CASE-BASED DECISION TREE

Risk	Failure Mode	Impact	Likelihood	Mitigation	Detection	Containment
Users skip questionnaire	Feature unused → No guidance	H	M	Make optional, mobile drawer, show value promise	Completion >40%	Disable; use translator only
Expert tagging wrong	Wrong specs → Trust loss	H	M	Validate with reviews, 3+ experts, pilot 3 categories, user feedback	Post-purchase: >70% yes	Pause scaling; re-interview
Logic fails	Wrong specs highlighted	H	M	Simple if-then rules, QA 20 scenarios, test 10 users, manual override	A/B test conversion lift	Pause; retest & simplify
Expert shortage	Tagging delayed/incomplete	M	H	Hire ASAP, 5+ yrs credential, start with 3 categories	Week 8: 3 experts done?	Extend 4 weeks, hire contractor

SOLUTION 2: PLAIN-ENGLISH SPEC TRANSLATOR

Risk	Failure Mode	Impact	Likelihood	Mitigation	Detection	Containment
Translation quality poor	Inaccurate/vague → Confused	H	M	Template format, expert review per 10, beta test 5 users	Clarity: >80%	Rewrite low-scoring
Review data bias	Skewed quotes → Misleading	M	M	Relevant quotes only, 3+ & 3- balanced, show count, manual vet	Survey: >75% helpful	Remove review feature
Writer capacity	Rush → Shallow translations	M	H	20-30/week per writer, 2-3 writers or 6-8 weeks, expert review	Velocity: >25/week	Extend timeline, hire 2nd writer
Tooltip UX breaks	Users don't find/use → Unused	M	M	Desktop hover, mobile tap drawer, nudge animation, test 3 devices	Usage: >30% tap	Show in-line instead

COMBINED RISKS (Both Solutions)

Risk	Failure Mode	Impact	Likelihood	Mitigation	Detection	Containment
No synergy	Separate features → No conversion lift	H	M	A/B test 4 combos: neither, tree only,	Both group >10% lift	Disable one; redesign flow

Risk	Failure Mode	Impact	Likelihood	Mitigation	Detection	Containment
				translator only, both		
Low trust	Users prefer YouTube → Ignored	H	M	Interview users, add expert credentials, link reviews, A/B test signals	Adoption: >40%	Add credibility signals

▼ Detailed Risk Description & Mitigation

SOLUTION 1: USE-CASE-BASED DECISION TREE - 5 KEY RISKS

Risk #1: Users Skip Questionnaire

Aspect	Details
Assumption	Users will engage with 3-4 minute questionnaire
Failure Mode	Users skip it → Feature unused → No guidance benefit
Impact/Likelihood	H / M
Mitigation	Make optional (appear after 2+ min on page), use low-friction mobile drawer, show value promise
Detection Signal	Completion rate >40% (if <20%, rollback)
Containment	Disable feature; rely on translator + importance scores

Risk #2: Expert Tagging is Wrong

Aspect	Details
Assumption	Experts correctly identify 3-5 decision-critical specs
Failure Mode	Subjective/wrong tagging → Users confused → Trust loss
Impact/Likelihood	H / M
Mitigation	Cross-check with review data (40%+ mentions = validated), use 3+ experts for consensus, pilot 3 categories first, allow user feedback to refine
Detection Signal	Post-purchase survey: >70% say specs helped (if <50%, revise)
Containment	Pause scaling; re-interview 20 customers per category; adjust criteria

Risk #3: Question-to-Spec Logic Fails

Aspect	Details
Assumption	Questions correctly map to specs (e.g., "bedroom" + "noise sensitive" = focus on Noise)
Failure Mode	Wrong specs highlighted → User frustration → Decision quality worsens
Impact/Likelihood	H / M
Mitigation	Use simple if-then logic (not ML), QA test 20 scenarios, user test with 10 target users, build manual override

Aspect	Details
Detection Signal	A/B test: Decision Tree group conversion (if lower than control, logic failed)
Containment	Pause; retest logic with QA and users; simplify questions

Risk #4: Explanations Unclear

Aspect	Details
Assumption	1-line specs ("CADR 300 = covers 300 sq ft") are clear
Failure Mode	Explanations too technical/vague → Users still confused → No time savings
Impact/Likelihood	M / M
Mitigation	Write for 8th grade reading level, test clarity with 5 users, include real-world context + cost implications, crowdsource refinement with feedback
Detection Signal	User rating: >75% say clear & helpful (if <60%, rewrite)
Containment	A/B test revised explanations; roll out better version

Risk #5: Expert Resource Shortage

Aspect	Details
Assumption	Experts available for 4 weeks of feature tagging
Failure Mode	Expert shortage → Incomplete tagging → Launch delayed or low-quality
Impact/Likelihood	M / H
Mitigation	Hire ASAP (not week 5), define credentials (5+ yrs or 100+ purchases), start with 3 categories only, create detailed tagging guide
Detection Signal	Did 3+ experts complete tagging for 3 categories by Week 8?
Containment	Extend timeline 4 weeks, hire contractor, outsource to consulting firm

SOLUTION 2: PLAIN-ENGLISH SPEC TRANSLATOR - 5 KEY RISKS

Risk #1: Translation Quality/Accuracy

Aspect	Details
Assumption	Writers translate 100+ specs accurately
Failure Mode	Inaccurate/vague translations → Users still confused
Impact/Likelihood	H / M
Mitigation	Use template: [SPEC] = [METRIC] = [CONTEXT] = [WHO CARES], expert reviews every 10, beta test with 5 users (>80% say clear)
Detection Signal	User rating: >80% clear & helpful (if <60%, rewrite)
Containment	Pause rollout; rewrite low-scoring translations; re-verify all

Risk #2: Review Data Bias

Aspect	Details
Assumption	Review quotes improve credibility without bias
Failure Mode	Skewed reviews (extreme opinions only) → Misleading users

Aspect	Details
Impact/Likelihood	M / M
Mitigation	Only surface quotes for that spec, require 3 positive AND 3 negative mentions (balanced), show quote count for transparency, manual vetting
Detection Signal	User survey: >75% found quotes helpful (if <60%, disable)
Containment	Remove review quote feature; rely on translation alone

Risk #3: Translation Inconsistency

Aspect	Details
Assumption	Same spec explained consistently across categories
Failure Mode	Same spec explained differently → Users confused → Trust loss
Impact/Likelihood	M / M
Mitigation	Create master glossary (1 translation per spec, reused), monthly audits across 3 random products, centralized management
Detection Signal	Audit: >95% specs consistent (if <85%, trigger rebuild)
Containment	Freeze new translations; rebuild glossary; re-launch

Risk #4: Tooltip UX Problems

Aspect	Details
Assumption	Hover/tap tooltips used naturally without friction
Failure Mode	Users don't notice/use tooltip; breaks on mobile → Feature unused
Impact/Likelihood	M / M
Mitigation	Desktop: hover (standard), Mobile: tap drawer (not overlay), first-time nudge with animation, test on 3 devices
Detection Signal	Usage rate: >30% tap at least once (if <15%, redesign)
Containment	Show translation in-line (always visible) instead of behind tap

Risk #5: Writer Capacity/Burnout

Aspect	Details
Assumption	Writers produce 100+ translations in 4 weeks
Failure Mode	Rush → shallow translations → users don't save time
Impact/Likelihood	M / H
Mitigation	Realistic scope: 20-30/week per writer, 2-3 writers or 6-8 weeks, expert review per 10, gold standard examples, weekly check-ins
Detection Signal	Velocity: >25 translations/week (if <15/week, rescope)
Containment	Extend timeline 6-8 weeks, hire additional writer, outsource to agency

COMBINED RISKS (Both Solutions) - 2 KEY RISKS

Combined Risk #1: No Synergy

Aspect	Details
Assumption	Decision Tree + Translator together improve conversion +10-12%
Failure Mode	Features work separately but don't combine → Users still take 1-2 hours, no conversion lift
Impact/Likelihood	H / M
Mitigation	A/B test all 4 combos: (1) Neither, (2) Tree only, (3) Translator only, (4) Both, set synergy bar at 10%+
Detection Signal	A/B test: Both group >10% lift (if <5%, synergy failed)
Containment	Pause one feature; double-down on the other; redesign UX flow

Combined Risk #2: Low Trust in Guidance

Aspect	Details
Assumption	Users trust Amazon guidance more than YouTube/ChatGPT/friends
Failure Mode	Users still prefer external sources → Guidance ignored → Goal not achieved
Impact/Likelihood	H / M
Mitigation	Interview 5 users post-launch why they did/didn't use, surface expert credentials on Decision Tree, link to reviews showing relevance, A/B test credibility signals
Detection Signal	Adoption: >40% use features (if <20%, trust issue)
Containment	Add credibility signals; if <30% adoption still, pivot to different solution

MONITORING DASHBOARD (Weekly During Pilot)

Feature	Metric	Target	Action if Missed
Decision Tree	Completion rate	>40%	If <20%, make optional
	Post-decision help	>70% yes	If <50%, revise tagging
	A/B conversion lift	+7-10%	If <3%, retest logic
	Research time	30 min	If >45 min, check UX
Translator	Clarity rating	>80%	If <60%, rewrite
	Tap usage	>30%	If <15%, redesign UX
	Review quotes trust	>75%	If <60%, disable
Combined	Synergy lift	+10-12%	If <5%, redesign flow
	Adoption rate	>40%	If <20%, add credibility
	Final conversion	+10-12%	If <7%, revisit strategy