

Lovely App! Don't Look Inside.



Are AI App Builders Secure?

We Tested  **Lovable**,  **Base44**, and **bolt.new** to Find Out.

Eran Cohen,
OX Security Research



Executive Summary

AI App Builders promise that anyone can build and deploy a production application in minutes. But our testing of Lovable, Base44, and Bolt proves that their generated code is insecure by default - even when prompting with the word "secure."

When we asked these three popular AI app builders to create a simple wiki app, all generated exploitable XSS vulnerabilities. Even when we explicitly asked for "secure" code, the vulnerabilities persisted. **Only detailed security instructions prevented the flaws. The platforms' built-in security scanners either missed the vulnerabilities or were inconsistent about their findings.**

This exposes a critical gap: AI builders marketed to non-technical users are not secure by default, and their security features create a dangerous false sense of safety.

Introduction: The AI App Builder Promise

"I love Lovable," Nvidia CEO Jensen Huang said in a recent CNBC interview, naming the AI app builder as one of the fastest-growing companies in enterprise AI. His enthusiasm is backed by extraordinary market momentum: Last July, Lovable raised \$200M at a \$1.8B valuation just eight months after launch. Base44, a solo-owned project just six months old, sold to Wix for \$80M cash in June. Bolt raised \$105.5M in a Series B round that valued the company at \$700M in January 2025.

These aren't just popular tools - they represent a massive bet on the future of software development. But do speed and ease come at a security cost?

AI app builders promise something extraordinary: anyone can build and deploy a production-ready application in minutes, no coding required. Just describe what you want, and moments later, a working app appears. "For most people, a computer science degree is no longer the entry ticket," said Lovable CEO Anton Osika in an interview with Business Insider last August. "You can build, ship, and even start companies without it."

Before you hang up your diploma for good, we wanted to ask a few uncomfortable questions: Do AI app builders generate secure code by default? Can their built-in safeguards catch vulnerabilities? And when things go wrong, who's accountable - the platform or the user?

To find out how secure these AI-generated apps really are, we put three leading platforms to the test - and found a critical gap: **When asked to create a simple wiki app - even when asking for security in the prompt - Lovable, Bolt, and Base44 all produced code with a classic stored XSS vulnerability, and their built-in security scans either missed it or flagged it inconsistently.**

The Test

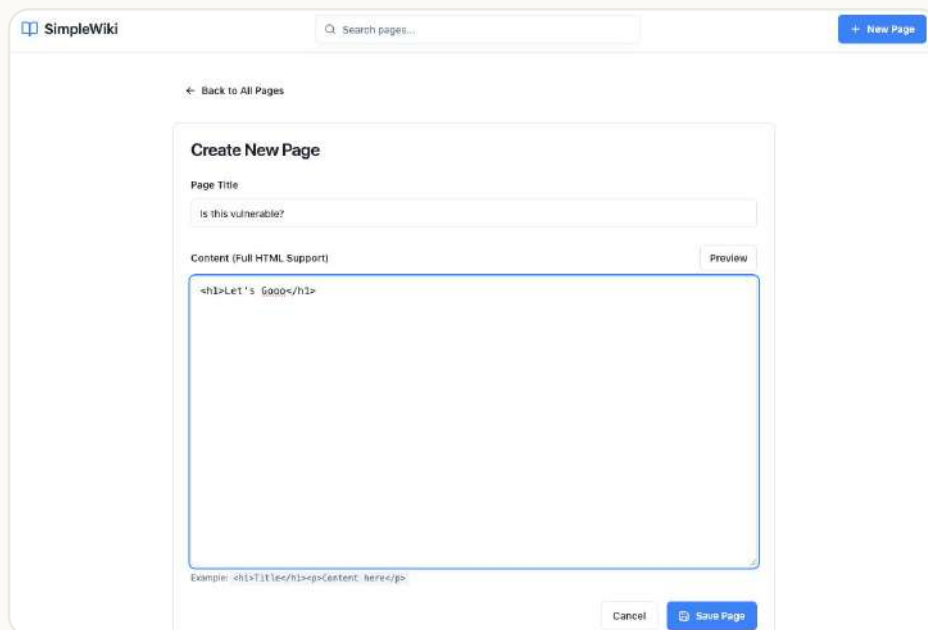
Phase 1: Testing the Baseline

Our testing focused on a fundamental web application security principle: **sanitizing user-generated content before it is stored and rendered.**

We asked one of the builders, Lovable (the methodology was consistent for all three platforms), to create a simple Wiki app:

"Create a wiki app - allow page creation with html tags editor support"

The resulting Wiki included a rich-text editor, which is where the security flaws emerged.



The vulnerable app we created through prompting (Lovable)

The Security Issue We Found - A Classic Stored XSS

The most basic defense against XSS is blocking common payload tags like **<script>**.

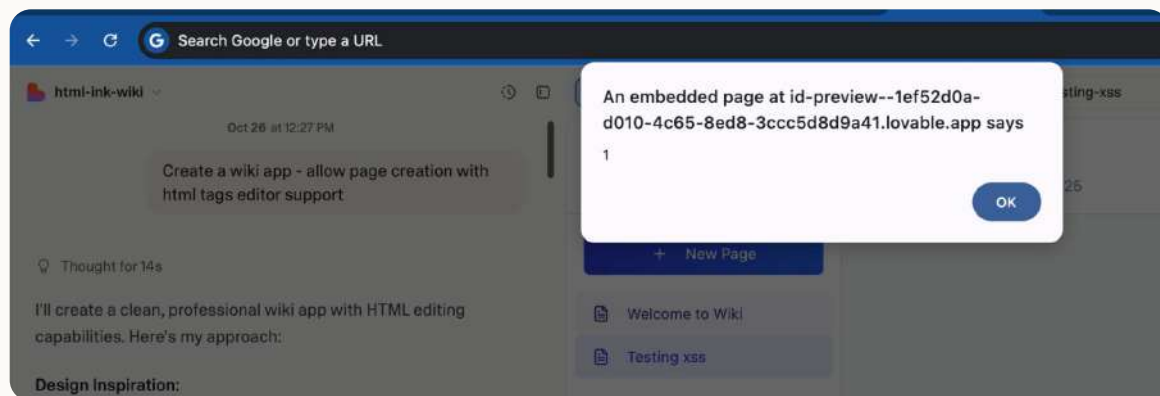
<script>alert(1)</script> - was successfully blocked or stripped *within the editor UI*.

However, a slightly less obvious, but equally effective, XSS vector slipped through in every tested platform:

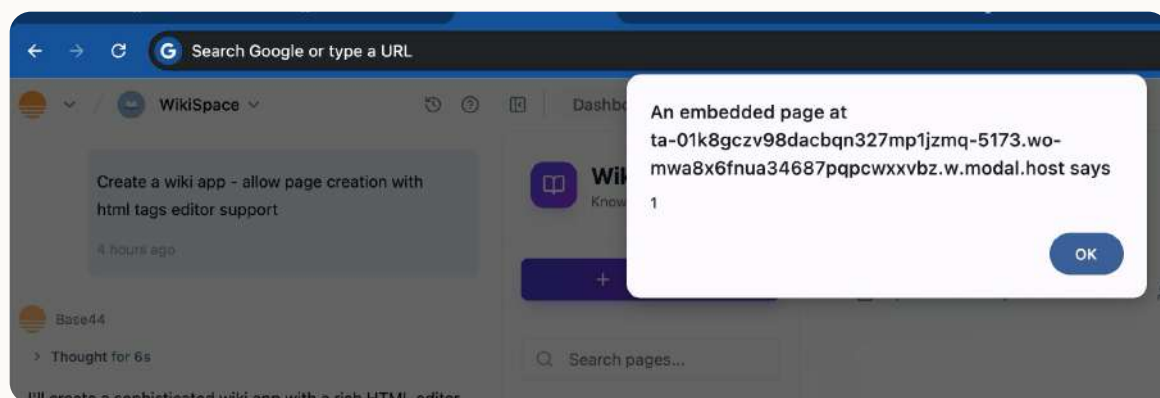
**** - was not blocked.

This payload exploits an HTML image tag's error handler to execute JavaScript. When the browser tries to load the non-existent image source (x), it triggers the **onerror** event, which executes the malicious code. This technique often bypasses simple XSS filters that only look for **<script>** tags.

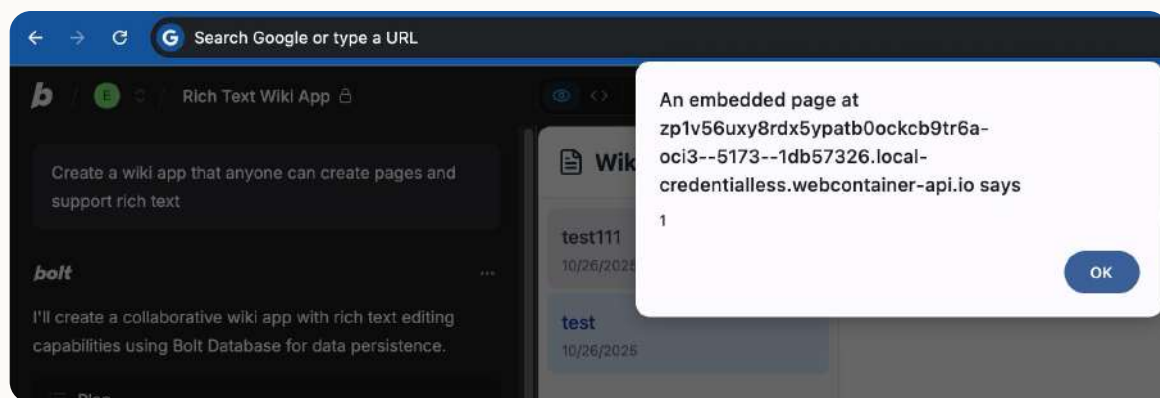
Lovable Stored XSS:



Base44 Stored XSS:



Bolt.new Stored XSS:



Potential damage: Session Hijacking and Data Theft

The risk of a Stored XSS goes far beyond a simple `alert(1)`. Applications built on these platforms often store sensitive user data, such as an **authentication token** (which grants access to that specific app instance), in the browser's **localStorage**.

An attacker can craft a payload to silently steal this token and send it to the attacker server:

```

```

This means a malicious user could inject this into a public page, and anyone who views it could have their session hijacked, leading to unauthorized access or data theft.

Here's a demonstration showing the stored XSS successfully fetching the victim's authentication token (on Base44 app):



The failure was not in the app's features - as requested by the user - but in the **underlying framework** logic provided by the AI App Builder itself, which is responsible for storing and rendering user-controlled HTML.

Phase 2: Testing Built-in Security Checks



Lovable, Bolt, Base44, and similar builders often feature some kind of "Built-in Security Check". Each platform takes a different approach to security scanning, but all share a fundamental flaw: **none of them prevent vulnerable code from being generated in the first place.**

Critically, none of these platforms require fixes to be applied before publishing. Security scanning functions as an optional afterthought rather than a fundamental design constraint.

Best Practice Would Require:

- 01 Generating secure code from the start
- 02 Enforcing secure sanitization libraries as non-negotiable defaults when the system accepts and renders user-supplied HTML
- 03 Making security an inherent design constraint, not an optional afterthought

When a platform markets itself to non-technical users and promises to handle the technical implementation securely, relying on post-generation scanning is insufficient.

Next: The Nuance: Where the Scanners Differ

The Nuance: Where the Scanners Differ



Platform	AI Generated Vulnerability	Scan Runs Before Deploy	Scan Requires Credits	Fix Requires Credits	Vulnerability Detected
 Lovable	Yes	Automatic	No	No	2 out of 3 attempts
 Base44	Yes	Manual only	No	No	No
 bolt.new	Yes	Automatic	No	Yes	No

The results of our test reveal significant differences in both the implementation and effectiveness of the security scanners - and critical gaps in what they're designed to detect in the first place.

Lovable's Inconsistent Detection:

Lovable automatically runs a security scan before publishing, which is a positive step. When the scan successfully detected the Stored XSS vulnerability (which happened in just 2 out of 3 attempts), it correctly suggested implementing DOMPurify to sanitize user input. After user approval, the fix was properly applied at no cost.

However, this inconsistency is deeply concerning - the same vulnerable code pattern was flagged in some generations but missed in others. **Users cannot rely on a scanner that catches the same vulnerability only 66% of the time.**

This inconsistency highlights a fundamental limitation of AI-powered security scanning: because AI models are non-deterministic by nature, they may produce different results for identical inputs. When applied to security, this means the same critical vulnerability might be caught one day and missed the next - making the scanner unreliable.

Inconsistent detection is arguably worse than no detection at all - as it creates false confidence while providing unreliable protection.

Base44's Manual-Only Approach:

Base44 provides a security scanning feature, but it's entirely manual - there's no automatic check before publishing, and users must actively remember to run it. In our testing, when we did manually trigger the scan, it failed to identify the XSS flaw in any attempt, giving the vulnerable app a "clean bill of health." **This represents the worst of both worlds: optional security that doesn't work when you do use it.**

Bolt's Financial Disincentive:

Bolt automatically runs a security scan before publishing, but charges credits to apply fixes. This creates a perverse incentive structure where users are informed about vulnerabilities but must pay to fix them. In our testing, **Bolt's scanner failed to detect the XSS vulnerability, but even if it had, the credit cost for fixes could discourage users from actually securing their applications.**

Base44 & Bolt: The Scope Problem

According to Base44 and Bolt's technical documentation, **their built-in security checks are designed to detect only basic issues such as missing access rules, unsafe backend function exposure, or secrets left in frontend code.** Many vulnerability classes - including XSS - fall completely outside their security check scope by design.

When users run the security check within the app interface, nothing indicates this narrow scope. Users expecting comprehensive security coverage must proactively hunt through platform documentation to discover these limitations.

This creates a dangerous gap between user expectations and actual protection. A non-technical user who sees "Security Check: Passed" has no reason to suspect that common, exploitable vulnerabilities like XSS were never even scanned for. The scanner isn't just failing to catch XSS - it was never designed to look for it, and users aren't told this upfront where it matters.

When platforms market themselves to non-technical users with promises of handling "security practices" (Base44's language), failing to clearly communicate the limitations of their security tools in the interface itself leaves users vulnerable.

Phase 3: Prompting for Security

The initial failure of the AI builders stemmed from a lack of secure defaults. We next tested whether a user could overcome this flaw by explicitly adding security requirements to the initial prompt.

We tested two prompts:

Test Prompt	AI App Builders Results
Test 1: "Create a Secured wiki app - allow page creation with html tags editor support."	ALL FAILED. Initial code still contained the XSS flaw.
Test 2: "Create a wiki app - allow page creation with html tags editor support (users should be able to use <h1>, etc.) Make sure to follow secure coding and preventing owasp top 10 issues"	Lovable - Success - Installed DOMPurify to sanitize user input and prevent XSS attacks. Base44 and Bolt - Partial success - Instead of using a vetted sanitizer (e.g. DOMPurify), they implemented an in-app XSS sanitizer. It prevented the basic XSS we tested but is not bullet-proof . Recommend switching to a proven library.

The takeaway is clear: The AI's security outcome is heavily influenced by the **specificity of the user's prompt**. Simply asking the AI for a "Secure app" is insufficient. While providing detailed security instruction can prompt the AI to introduce defenses, users must still **manually verify** that the AI chose a secure, proven library (like **DOMPurify**) and did not rely on unvetted, custom security code.

Unfortunately, this means that far from the promise - non-technical users won't know what to ask for.

Platform	Basic Prompt	"Secured" Prompt	Explicit OWASP/Security Instructions
 Lovable	Failed (XSS present)	Failed (XSS present)	Success: DOMPurify implemented
 Base44	Failed (XSS present)	Failed (XSS present)	Partial: Custom sanitizer (not vetted)
bolt.new	Failed (XSS present)	Failed (XSS present)	Partial: Custom sanitizer (not vetted)

Lovable Test #2 Result:

Security measures:

- DOMPurify for HTML sanitization
- Input validation with length limits
- No dangerous HTML attributes/tags
- Secure rendering practices

Let me implement this:

- Installed dompurify@latest
- Installed @types/dompurify...
- Edited index.css
- Edited tailwind.config.ts
- Edited wiki.ts
- Edited sanitize.ts
- Edited wikiStorage.ts
- Edited WikiLayout.tsx
- Edited WikiEditor.tsx
- Edited WikiContent.tsx
- Edited WikiHome.tsx

Back to Preview

Code

Ask Lovable...

Files

Search

6 results in 2 files

- package.json


```
43 "types/dompurify": "^3.3.0",
48 "dompurify": "^3.3.0",
```
- src/lib/sanitize.ts


```
1 import DOMPurify from 'dompurify';
1 import DOMPurify from 'dompurify';
3 // Configure DOMPurify to allow safe HTML tags while preventing XSS
34 return DOMPurify.sanitize(html, sanitizeConfig);
```

src/pages/Index.tsx

```
1 import DOMPurify from 'dompurify';
2
3 // Configure DOMPurify to allow safe HTML tags while preventing XSS
4 const sanitizeConfig = {
5   ALLOWED_TAGS: [
6     'h1', 'h2', 'h3', 'h4', 'h5', 'h6',
7     'p', 'br', 'hr',
8     'strong', 'em', 'u', 's', 'code', 'pre',
9     'ul', 'ol', 'li',
10    'a', 'img',
11    'blockquote',
12    'table', 'thead', 'tbody', 'tr', 'th', 'td',
13    'div', 'span'
14  ],
15   ALLOWED_ATTR: [
16     'href', 'src', 'alt', 'title', 'class'
17   ],
18   ALLOWED_URI_REGEX: /^(?:https?:|mailto:|tel:|data:|[\^a-z
19 ]+)$/;
20
21 /**
22  * Sanitizes HTML content to prevent XSS attacks
23  * Follows OWASP guidelines for secure HTML handling
24  */
25 export const sanitizeHTML = (html: string): string => {
26   if (!html || typeof html !== 'string') {
27     return '';
28   }
29
30   // Limit content length to prevent DoS
31   const maxLength = 100000;
32   const truncatedHtml = html.slice(0, maxLength);
33
34   return DOMPurify.sanitize(truncatedHtml, sanitizeConfig);
35 }
```

Base44 Test #2 Result:

Back to Dashboard

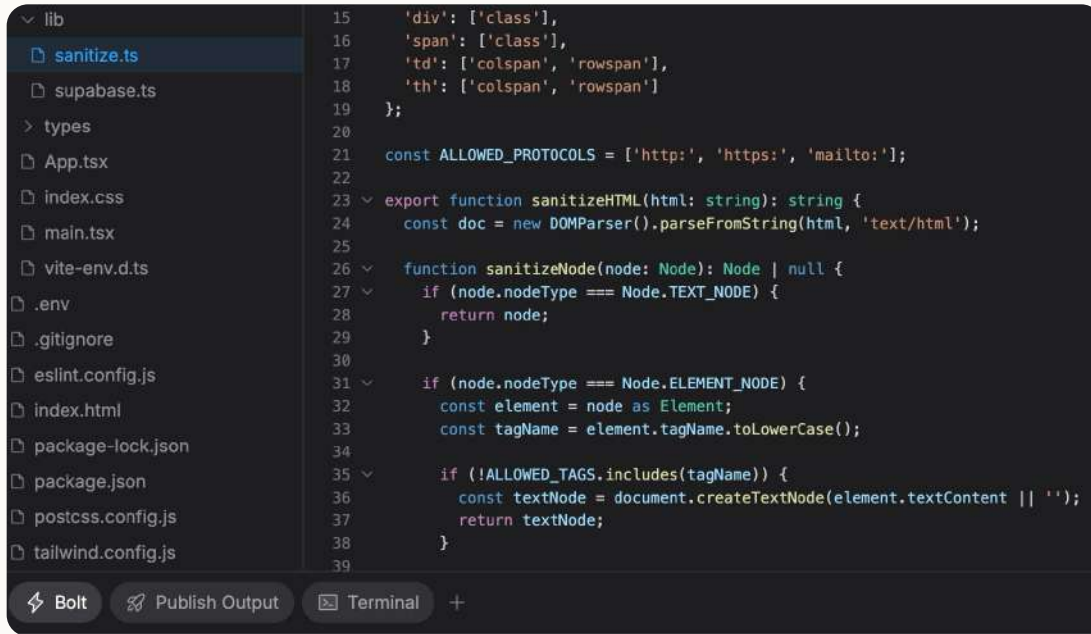
Search in files...

- Pages
- Components
 - wiki
 - RichTextEditor
 - WikiPageCard
 - HtmlSanitizer
- Entities
- Layout.js

Code editing is only available on paid plans. [Upgrade your plan](#)

```
1 /**
2  * HTML Sanitizer to prevent XSS attacks (OWASP Top 10)
3  * Implements a whitelist approach for safe HTML rendering
4  */
5
6  const ALLOWED_TAGS = [
7    'h1', 'h2', 'h3', 'h4', 'h5', 'h6',
8    'p', 'br', 'hr',
9    'strong', 'b', 'em', 'i', 'u', 's', 'mark',
10   'ul', 'ol', 'li',
11   'a', 'img',
12   'blockquote', 'pre', 'code',
13   'table', 'thead', 'tbody', 'tr', 'th', 'td',
14   'div', 'span',
15   'video', 'audio', 'source',
16   'dl', 'dt', 'dd',
17   'abbr', 'cite', 'kbd', 'sub', 'sup'
18 ];
19
20 const ALLOWED_ATTRIBUTES = {
21   'a': ['href', 'title', 'target', 'rel'],
22   'img': ['src', 'alt', 'title', 'width', 'height'],
23   'video': ['src', 'width', 'height', 'controls', 'poster'],
24   'audio': ['src', 'controls'],
25   'source': ['src', 'type'],
26   'td': ['colspan', 'rowspan'],
27   'th': ['colspan', 'rowspan'],
28   '*': ['class', 'id', 'style'] // Global attributes (limited)
29 };
30
31 const ALLOWED_PROTOCOLS = ['http:', 'https:', 'mailto:', 'tel:'];
32
33 /**
34  * Sanitizes HTML content by removing dangerous elements and attributes
35  * @param {string} html - Raw HTML string
36  * @returns {string} - Sanitized HTML string
37  */
38 export function sanitizeHtml(html) {
39   if (!html) return '';
40 }
```

Bolt Test #2 Result:



```

15   'div': ['class'],
16   'span': ['class'],
17   'td': ['colspan', 'rowspan'],
18   'th': ['colspan', 'rowspan']
19   };
20
21   const ALLOWED_PROTOCOLS = ['http:', 'https:', 'mailto:'];
22
23   export function sanitizeHTML(html: string): string {
24     const doc = new DOMParser().parseFromString(html, 'text/html');
25
26     function sanitizeNode(node: Node): Node | null {
27       if (node.nodeType === Node.TEXT_NODE) {
28         return node;
29       }
30
31       if (node.nodeType === Node.ELEMENT_NODE) {
32         const element = node as Element;
33         const tagName = element.tagName.toLowerCase();
34
35         if (!ALLOWED_TAGS.includes(tagName)) {
36           const textNode = document.createTextNode(element.textContent || '');
37           return textNode;
38         }
39

```

The Accountability Gap: Who's Responsible?

When we reported this issue, **Base44** replied:

"After review, this issue is not considered a vulnerability in Base44 itself. The XSS occurs because of how the custom Wiki app was created, not because of a flaw in the Base44 framework."

Lovable response:

"After review, there doesn't seem to be any significant security impact as a result of the behavior you are describing. This is an application generated by Lovable, any misconfiguration on this is not part of our program Scope."

Technically, these vendors might have a point - the payload lived in an app generated by their AI based on user input.

But this response directly contradicts their own marketing promises. On Base44 website, for example - they explicitly state: "Our platform is designed to be easily accessible to non-technical users. Just describe your software needs in plain language, and our AI will handle the technical implementation." They further claim: "we take data security very seriously. User management and authentication systems are built-in, using best-in-class, industry-standard encryption and security practices to protect your data and your users' information."

Generating code with trivially exploitable XSS vulnerabilities is a failure to deliver on those promises.

The user didn't write `dangerouslySetInnerHTML` or choose to skip `DOMPurify` - the AI did. When a platform markets itself to non-technical users and promises to handle the technical implementation securely, it should not then claim that security vulnerabilities in that AI-generated implementation are out of scope.

These platforms position themselves as end-to-end solutions that abstract away the complexity of application development. They should enforce secure defaults in the framework code they automatically generate, especially for well-understood vulnerability classes like XSS. When the AI is the architect, the developer, and the security reviewer, the platform bears responsibility for the security outcomes.

Bolt was contacted using the same disclosure process, but had not responded by time of publication.

Key Takeaways for Users Using AI App Builders

AI will happily build features for you, but it won't tell you which ones are exploitable. This is a critical lesson for users of **Lovable**, **Bolt**, **Base44**, and all similar tools.

For AI Builders Users:

- **Expect AI builders to generate code containing software vulnerabilities**
- **Explicitly prompt for security: and be specific about it.** Force the AI to use secure patterns by adding directives like: "All user-generated HTML must be strictly sanitized using DOMPurify."
- **Run your own Security tools:** Implement tools like **SAST**, **DAST** and **WAF** on your AI-generated code. AI should complement, not replace, traditional security tools and manual code reviews.
- **Sanitize all the things:** Manually ensure that all user-generated content is sanitized using a robust, community-vetted library like **DOMPurify** (for client-side) or a reliable server-side library.

Conclusion: Creating the Perfect Storm, One Prompt at a Time

In our recent **"Army of Juniors"** report, OX identified two critical effects of AI-generated code: "Insecure by Dumbness" - when non-technical users deploy applications without understanding the security implications. And **"The It Works Trap"** - when functional code passes basic tests but harbors serious vulnerabilities beneath the surface.

AI app builders like Lovable, Bolt, and Base44 represent the most extreme manifestation of both phenomena. As our testing reveals, surface functionality masks fundamental security flaws that these platforms neither prevent during generation nor reliably detect in their security scans.

When the AI is both the junior developer and the security reviewer, and the user lacks the expertise to question either, we've created the perfect storm for insecure-by-default applications at scale.

About OX

OX Security is a leader in application and product security, providing the most comprehensive coverage in the industry throughout the entire software development lifecycle, from code to runtime through the cloud.

OX created **VibeSec, the first Vibe Security platform** that prevents insecure AI-generated code before it exists, embedding real-time security context directly into AI coding agents at the moment of creation.

[Learn More](#)