

When AI Misleads: Singapore Courts Signal Stricter Liability for Unverified AI Output

(Summary and analysis by ABU Legal Division)

Summary of Recent Legal Developments

Courts in both the United States and Singapore are moving decisively toward treating AI tools not as neutral platforms but as commercial products capable of causing harm.

In the US, a high-profile lawsuit was allowed to proceed in 2025 involving an AI chatbot that allegedly contributed to a teenager's death. The plaintiffs argue that the chatbot was defectively designed and that the developer should be held liable in the same way manufacturers are responsible for unsafe products. This represents a significant shift, from "AI as software" to "AI as a commercial product."

In Singapore, the High Court delivered a strong signal in September 2025. A lawyer was sanctioned for submitting fabricated citations generated by an AI tool. The court held that failing to verify AI output amounts to negligence, and that professionals have a personal duty to check the accuracy of AI-generated information.

Together, these developments show that courts are increasingly unwilling to accept the argument that AI errors are unforeseeable or excusable. Businesses that deploy AI in customer-facing roles may be held directly accountable for misleading or harmful statements generated by their systems.

Key Legal Risks for Businesses (Including Broadcasters)

1. Negligence: Duty to Verify AI Output

Singapore courts have made it clear that relying on unverified AI output, especially in a professional or commercial context, will be viewed as negligent.

- a. A business cannot hide behind disclaimers such as "AI-generated content may contain errors."
- b. If an AI system provides misleading information to viewers, advertisers, or subscribers, liability will very likely fall on the organisation, not the vendor.

2. Expectation and Reliance

If an AI system communicates with consumers in a way that resembles customer service, courts will consider that communication binding.

- a. For example, an AI bot that promises a refund or confirms a subscription change creates a "reasonable expectation" that must be honoured.
- b. Fine-print disclaimers will not protect a business if the overall service delivered is substantially different from what the user was led to expect.

3. No Safe Harbour for Generative AI in Singapore

Unlike the US (which has Section 230-style protections), Singapore does not offer broad immunity for platform content.

- a. Businesses may be treated as "publishers" of AI-created messages.
- b. If a corporate chatbot defames a competitor or misleads a member of the public, the organisation may be held directly responsible.

4. **Third-Party AI Vendors and Supply Chain Liability**

Businesses that license AI systems face a “liability sandwich”:

- a. They are accountable to consumers downstream,
- b. Yet dependent on vendors upstream.
Updated procurement contracts must therefore include:
- c. **Warranties on system accuracy and compliance,**
- d. **Indemnities** for hallucinations, unsafe outputs, and design defects.

Implications for Broadcasters and Media Organisations

AI systems are increasingly embedded in operational workflows — from chatbots and viewer hotlines to newsroom automation, transcription, translation, scheduling, and content recommendations. These legal shifts mean:

- **AI is not treated as a tool; it is treated as a commercial agent.**
If an AI assistant makes a factual error in a financial broadcast or incorrectly moderates user comments, the broadcaster could be held liable.
- **Newsrooms must verify and not rely on AI output.**
Courts expect a “human in the loop,” especially for high-risk information such as legal, medical, election, or emergency-related content.
- **Disclaimers are not enough.**
Broadcasters need operational safeguards, documentation of verification processes, and transparent AI disclosures.

Recommendations for ABU Members

1. **Introduce clear AI disclosures**
 - a. Inform users when they are interacting with an AI system.
 - b. Explain the system’s limitations and known risks.
2. **Strengthen human oversight**
 - a. High-risk queries should be escalated to a human operator.
 - b. Newsroom AI output (summaries, headlines, translations) must be reviewed by editorial staff.
3. **Adopt a “verification-first” policy**
 - a. AI-generated claims, facts, and citations should never be published without human verification.
 - b. Document verification workflows to demonstrate due diligence.
4. **Revise procurement and vendor contracts**
 - a. Require AI vendors to provide warranties on accuracy and safety.
 - b. Include indemnities for model hallucinations, erroneous statements, and system failures.

For the full legal decision, read here: https://www.elitigation.sg/gd/s/2025_SGHCR_33