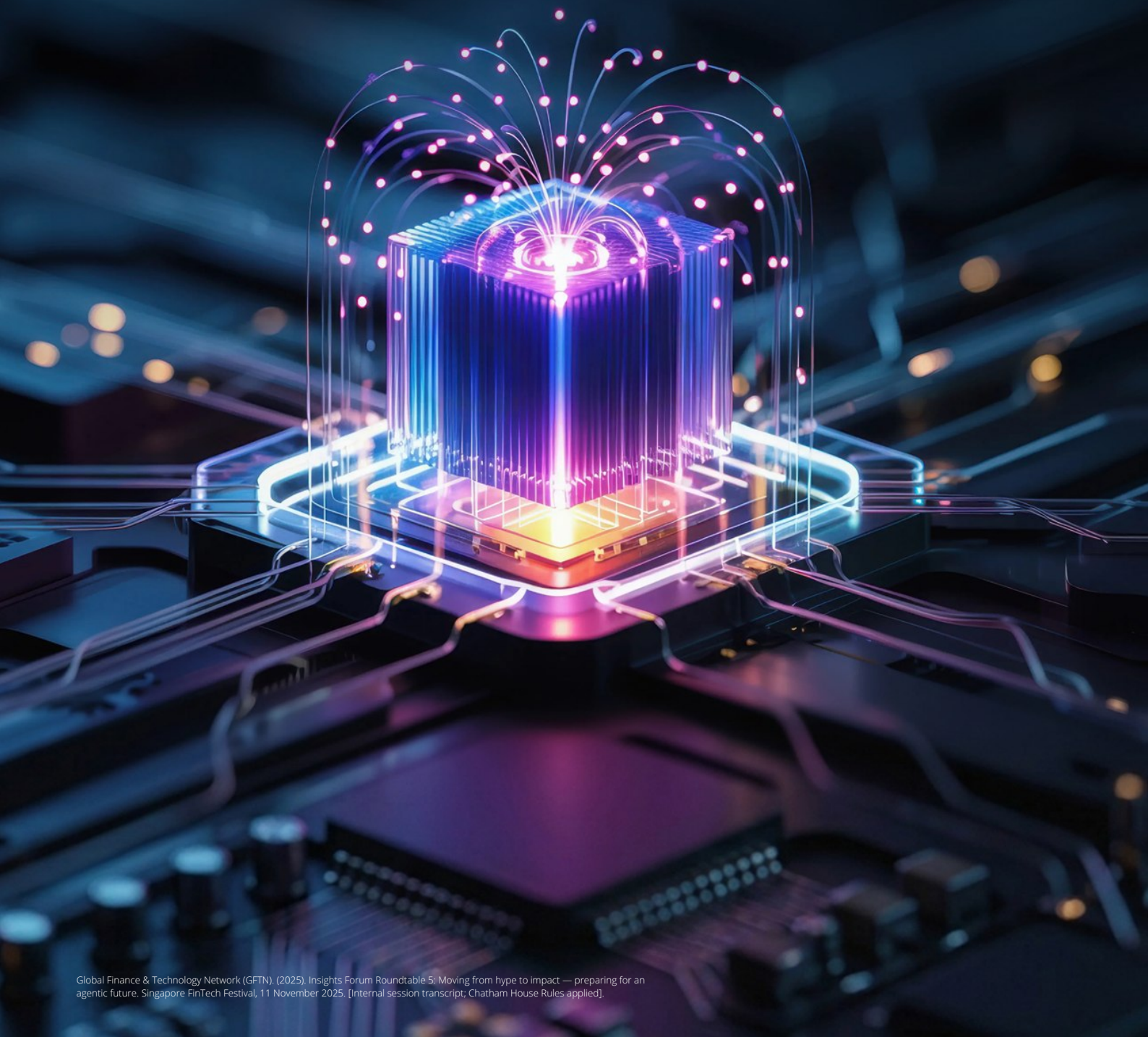


MOVING FROM HYPE TO IMPACT: PREPARING FOR AN AGENTIC FUTURE

June 2026



ABOUT

The Global Finance & Technology Network (GFTN) is a Singapore-headquartered organisation that leverages technology and innovation to create more efficient, resilient, and inclusive financial systems through global collaboration. GFTN hosts a worldwide network of forums (including its flagship event, the Singapore FinTech Festival); advises governments and companies on policies and the development of digital ecosystems and innovation within the financial sector; offers digital infrastructure solutions; and plans to invest in financial technology startups through its upcoming venture fund, with a focus on inclusion and sustainability.

For more information, visit www.gftn.co



Zühlke is a global transformation partner, with engineering and innovation in our DNA. We're trusted to help clients envision and build their businesses for the future – to run smarter today while adapting for tomorrow's markets, customers, and communities.

Our multidisciplinary teams specialise in tech strategy and business innovation, digital solutions and applications, and device and systems engineering. We excel in complex, regulated spaces including health and finance, connecting strategy, tech implementation, and operational services to help clients become more effective, resilient businesses.

Founded in Switzerland in 1968, Zühlke is owned by its partners and located across Europe and Asia. Our venture capital arm, Zühlke Ventures, provides start-up financing in HealthTech.

For more information, visit www.zuhlke.com



GFTN Insights is a year-round series that connects senior officials and industry leaders across continents to address key financial and tech challenges. Held under Chatham House rules, these dialogues foster partnership, culminating in the two-day Insights Forum in Singapore.

For more information, visit <https://gftn.co/global-forums/insights-forum>



EXECUTIVE SUMMARY

Agentic AI — the capacity of AI systems to autonomously plan, decide, and act across multi-step workflows — is no longer a future prospect for financial services. It is already here, unevenly deployed, frequently misunderstood, and almost universally undergoverned. Senior leaders from global banks, regulators, technology firms, and multilateral institutions gathered for an invite-only Insights Forum roundtable to examine where the industry truly stands.

The consensus was striking: the question has shifted from whether to adopt agentic AI to how to do so responsibly, at scale, and in a way that generates sustainable value. Yet, a persistent gap comes between ambition and execution. Across the table, participants agreed that the industry's dominant challenge has moved from technological to institutional. Fragmented governance, siloed data, misaligned expectations at the executive level, and a near-total absence of interoperability standards are stalling progress far more than model capability ever has.

This report distills five interconnected insights from that conversation, grounded in research and practitioner experience, with actionable pathways for financial institutions navigating the agentic transition.



of companies plan to put agents into production — yet only 11% have done so

KPMG, 2025



of agentic AI experiments are projected to fail by 2027 due to governance and coordination gaps

Gartner, 2025



CONTENTS

Executive Summary	03
1. Introduction	05
2. Insight 1: The real bottleneck is the data beneath the model	05
3. Insight 2: Human-in-the-loop is not sufficient, and over-reliance is its own risk	06
4. Insight 3: Governance must be designed in from the start instead of retrofitted	06
5. Insight 4: The back office is the killer use case, while most institutions are still looking at the front	07
6. Insight 5: Standards and talent are the limiting factors — in that order	08
7. Conclusion: Five imperatives for the agentic transition	08
8. References	09
Contributors	10

01 INTRODUCTION

Walk the exhibition floor of any major fintech event and two topics will find you before you find a seat: stablecoins and agentic AI. The pairing is telling. Both represent infrastructure-level bets: long-term commitments where early positioning determines future relevance. But where stablecoin adoption hinges primarily on regulatory clarity, agentic AI adoption is being held back by something altogether more tractable: a failure of institutional readiness.

Between the 2024 and 2025 editions of the Singapore FinTech Festival, the technical landscape shifted dramatically. The Model Context Protocol (MCP) went from niche developer discussion to industry standard. Agent-to-agent (A2A) communication frameworks emerged. Frontier models reached the capability threshold required for autonomous, multi-step task execution in enterprise environments. In short, the technology caught up. The institutions, largely, have not.

This is not a counsel of despair. The roundtable surfaced a cohort of practitioners who have moved well beyond the proof-of-concept phase — running hundreds of Gen AI use cases in production, deploying agentic systems in compliance, document processing, KYC, and capital markets analysis, and developing governance frameworks fit for autonomous systems. Their experience forms the empirical backbone of this report.

What follows is a synthesis of the structural insights it revealed, exploring five fault lines that will determine which institutions capture the value of agentic AI and which are left managing its risks without its rewards.



02 INSIGHT 1

The real bottleneck is the data beneath the model

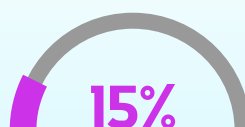
One of the most clarifying moments in any agentic AI deployment comes when a team realises it cannot access its own data. Vendor lock-in — a term that surfaced repeatedly across the roundtable — describes a less-than-ideal situation still happening today: financial institutions paying substantial fees to professional services firms simply to extract data from enterprise software they already own. As one participant from the RegTech space put it plainly:

“When you talk about vendor lock-in, you’re talking about how hard it is for financial services to get their own data out of their own systems. That should simply not be happening.”

The consequences of this structural fragility for agentic AI are significant. Agents require clean, accessible, well-governed data to function reliably. Without API-native architecture and interoperable data standards, organisations build isolated use cases rather than reusable data products — raising costs, slowing delivery, and preventing value from compounding across the ecosystem. The roundtable identified data accessibility as the single most cited regulatory constraint to AI adoption, with one central bank’s survey data showing that 33% of regulated firms identify challenges in data regulation as an active barrier.

The issue is compounded by the fundamentally text-centric and English-language bias of current agentic models. A senior participant from a major Middle Eastern bank described months of effort to achieve reliable Arabic-language output from agents. The challenge exposed a deeper gap: agentic AI, at its current frontier, reflects the data on which frontier models were trained. For institutions operating in multilingual, multi-jurisdictional contexts, this is a structural constraint on the scope of viable deployment.

The actionable implication is architectural. Institutions must treat API-native infrastructure as a governance precondition for agentic AI. Regulators, in turn, have an opportunity to accelerate progress here — not by introducing new AI-specific rules, but by establishing clear data portability and interoperability standards. As one participant from the regulatory community noted, guidance on data standardisation would address more of the practical barriers to responsible AI adoption than technology-specific rules ever could.



productivity loss faced by
companies without AI-
ready data foundations

IDC, 2025

03 INSIGHT 2

Human-in-the-loop is not sufficient, and over-reliance is its own risk

Every institution represented at the roundtable has a version of the same policy: no agentic output goes to a customer or regulator without a human reviewing it. This is sensible. But a striking observation from one participant running a banking operations platform revealed an uncomfortable truth: human oversight degrades when agents become too reliable.

"We got to a point where our operators were not reviewing submissions properly because they were so convinced of the reliability of the agentic output. We now have to deliberately insert mistakes so that reviewers stay sharp."

This is a governance failure of a novel kind: not the AI system failing, but the human oversight mechanism atrophying in response to AI competence. It points to a deeper design challenge in which human-in-the-loop is an architecture, not just a policy. It requires deliberate friction, periodic audits, and active calibration of reviewer attention. Without these, the label "human in the loop" becomes compliance theatre rather than meaningful oversight.

A separate but related tension surfaced in the discussion of speed. In capital markets and financial crime contexts, latency matters. A human review step that takes hours is not compatible with real-time transaction monitoring. This drove several participants to explore LLM-as-judge frameworks — where one model evaluates the output of another — as a way of maintaining oversight without sacrificing speed. One investment bank described a "reflection architecture" pairing a writer model and a reviewer model, achieving 70–80% quality output versus roughly 20% from a single zero-shot prompt.

The practical implication is a spectrum of operating models rather than a binary. Institutions must map each use case to the right point on a continuum from 'human in the loop' (full review) to 'human on the loop' (exception-based monitoring) to 'LLM as judge' (automated quality control). The choice should be driven by risk materiality, error tolerance, and time sensitivity — not by institutional habit or regulatory uncertainty.



of AI applications in UK-regulated firms were fully autonomous with no human review, as of 2024 — indicating not just policy, but genuine caution

Bank of England Survey, 2024

04 INSIGHT 3

Governance must be designed in from the start instead of retrofitted

If there is one statement that deserves to be carved into the wall of every AI strategy team, it is this: more autonomy equals more risk. The industry has long been wired to optimise for outcomes, and to treat the steps in between as implementation overhead. One participant put it plainly: in agentic AI, that instinct becomes a liability. Unlike conventional software, agentic systems require scrutiny at every step in the process, not just at the end. While that discipline is manageable in isolated, low-stakes applications, in multi-agent systems with real-world consequences, such as credit decisions, compliance filings, or customer communications, this is a liability.

"The failures we saw in multi-agent systems were not due to weak models. They were due to coordination breakdown — communication gaps, memory drift, cascading errors where a small mistake leaps across agents."

The roundtable surfaced a taxonomy of governance failures specific to agentic systems: coordination breakdown between agents; memory drift across long task chains; cascading errors, where early mistakes propagate and amplify; and loss of observability, where no single point in the system can account for a decision's full reasoning chain. These are not hypothetical risks. Multiple practitioners described them as live failures in production environments.

A critical structural gap identified by participants is the absence of 'compliance by design' thinking. Traditional software development centres the outcome; agentic development must centre the process, building in traceability, audit trails, and kill-switch mechanisms from the first line of code. The concept of 'state management and time travel' — the ability to rewind a workflow to any prior state in order to understand or correct failures — emerged as a practical design requirement, not a theoretical nicety.

Participants from the regulatory community offered an important calibration: AI has thus far fit reasonably well within existing risk management frameworks — model risk management policies, senior management accountability regimes, third-party risk rules. But the more autonomous end of the spectrum introduces features — self-evolving models, agent-to-agent delegation, cross-jurisdictional data flows — that will require frameworks to evolve. The appetite for proactive guidance appears to be growing on both sides.



of business leaders prioritise security, compliance, and auditability as their top concern when scaling agentic AI

KPMG, 2025

The Governance Design Checklist

Before deploying any agentic system:

Define kill-switch and pause mechanisms at design stage.

Require full chain-of-reasoning traceability as a system output, not an afterthought.

Map the use case to the appropriate oversight model (human in / on the loop or LLM-as-judge).

Integrate agentic risk into the enterprise risk taxonomy — including third-party, operational, and reputational dimensions.

Set explicit state management requirements to enable rollback and incident investigation.



05 INSIGHT 4

The back office is the killer use case, while most institutions are still looking at the front

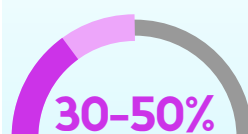
There is a persistent bias in how the financial services industry narrates its AI ambitions: customer-facing applications attract the headlines, the pilot budgets, and the board attention. Yet the roundtable participants who had moved furthest from experimentation to production had done so almost exclusively in the middle and back office, for a straightforward reason.

Banks carry some of the largest people cost bases of any industry. The majority of those people are performing middle and back office tasks, including reporting, compliance, reconciliation, onboarding, KYC, document processing. These tasks are characterised by high volume, defined rules, and low tolerance for error. These are precisely the conditions under which agentic AI delivers reliable, measurable results. As one participant operating five banks end-to-end observed, the industry is heading for “probably the biggest transfer in history from labour budgets to software budgets.”

“Banks spend around 60–70% of their costs on people. Those people are doing middle and back office activities. And those activities are remarkably well suited to agentic solutions.”

The concrete examples were instructive. One participant described an agentic system for Unusual Activity Monitoring (UAM) reporting that reduced a six-hour manual task to three to five minutes, with a ten-minute human review before submission to the regulator. Another described a KYC source-of-wealth assistant that supports relationship managers in synthesising complex client documentation to regulatory standards. A third outlined a compliance auto-review agent handling checklist validation across large document sets.

The common architecture across these successes is not, in fact, multi-agent orchestration. It is a single, well-scoped agent with a defined task, access to the right data, a clear human handoff point, and observability baked in. The roundtable reached a quiet consensus: attempting to build a fully autonomous, end-to-end agentic system from day one is a near-guarantee of failure. The path to that system runs through a sequence of small, production-grade wins.



reduction in manual workloads achieved by early agentic AI use cases in banking operations

McKinsey, 2025

06 INSIGHT 5

Standards and talent are the limiting factors — in that order

A consistent theme across the roundtable was sequencing: before organisations can build reliable agentic applications, the plumbing must work. That means interoperability standards at the agent and model level, and the workforce literacy to use them effectively. The discussion validated this ordering, and revealed why inverting it is so costly.

On standards: the absence of interoperability norms at the model and agent level is not a minor inconvenience. It is a structural inhibitor of scale. As one moderator noted, prompts, outputs, and agent behaviours do not transfer cleanly across model families, and mixing model lineages in multi-agent pipelines creates inconsistency at best and failure at worst. The emergence of the Model Context Protocol and A2A communication standards in the past year represents genuine progress, but industry-wide adoption remains nascent. Without agreed protocols for how agents identify themselves, delegate to other agents, and maintain accountability chains, multi-agent systems will continue to fail at the seams.

“When we were at SFF last year, there was no such thing as a Model Context Protocol or agent-to-agent communication framework. These emerged in the last six to seven months. The pace of change demands that governance keeps up.”

On talent: the roundtable surfaced a more nuanced concern than the familiar “we need more AI engineers” narrative. The critical gap is no longer technical headcount, but executive-level literacy. Multiple participants described the same dynamic: C-suite leaders holding AI teams to unrealistic timelines, creating pressure that shortcuts the governance and data work that determines whether deployments succeed or fail.

The implication is that upskilling programmes which focus only on technical practitioners are solving the wrong problem. The more urgent intervention is an executive education effort — one that builds a shared, accurate mental model of what agentic AI can and cannot do, on what timelines, under what conditions. Until boards and senior leadership teams have that literacy, the governance structures required for responsible agentic deployment will continue to be underfunded and under-prioritised.

A participant from the responsible AI space offered a practical model: begin with an organisational “pledge” — a foundational commitment to responsible AI practice that creates shared language and accountability across hierarchy levels. Layer on a governance body, a responsible AI framework, workforce development integration, and an incident reporting cadence. This structured approach, applied across industries, has demonstrably closed the gap between AI ambition and AI readiness.



of organisations are actively scaling agentic AI in at least one business function — while 39% remain in the experimentation phase

McKinsey, 2025

07 CONCLUSION

Five imperatives for the agentic transition

The roundtable ended with a series of wish lists, and the consistency across them was itself an insight. Participants did not wish for better models, faster chips, or cheaper compute. They wished for unlocked return on investment, mature agent-to-agent ecosystems with shared standards, responsible AI practice, broader upskilling, and the organisational patience to let good deployments mature. These are institutional desires, not technical ones.

For many use cases, the capability threshold has been crossed, with technology advancing faster than the institutions deploying it. The question is whether financial institutions, regulators, and technology partners can build the institutional layer that agentic AI requires to fulfil its potential. Based on the discussion at the Singapore FinTech Festival 2025, we close with five imperatives for any institution serious about the transition:

- 1. Treat API-native data architecture as a governance precondition, not a technology upgrade.** Agentic AI cannot function reliably on data it cannot access. Interoperability is the foundation.
- 2. Design governance in from the start.** Map every use case to the appropriate oversight model. Build kill-switches, audit trails, and rollback capabilities into the system architecture before the first line of agent code is written.
- 3. Start in the back office.** High-volume, rules-driven, internal workflows offer the best conditions for building production-grade agentic competence. The front-office ambitions can follow.
- 4. Invest in executive AI literacy before technical headcount.** Unrealistic expectations set at the leadership level are a more common cause of AI programme failure than model limitations.
- 5. Engage proactively on standards.** The institutions and regulators that shape interoperability norms will have an outsized influence on the architecture of the agentic ecosystem. Participation in standards bodies is a competitive act, not a philanthropic one.

“The conversation is no longer about whether to adopt agentic AI. It is about how to do so in a way that is governed, trusted, and genuinely value-creating for the institutions, customers, and economies we serve.”

08

REFERENCES

1. KPMG. (2025). AI Quarterly Pulse Survey.
KPMG LLP, April 2025.
<https://kpmg.com/us/en/articles/2025/ai-quarterly-pulse-survey.html>
2. Gartner. (2025). Gartner predicts over 40% of agentic AI projects will be canceled by end of 2027.
Gartner Press Release, 25 June 2025.
<https://www.gartner.com/en/newsroom/press-releases/2025-06-25-gartner-predicts-over-40-percent-of-agentic-ai-projects-will-be-canceled-by-end-of-2027>
3. IDC. (2025). FutureScape 2026: Worldwide Artificial Intelligence and Automation Predictions.
International Data Corporation, October 2025.
<https://www.businesswire.com/news/home/20251023490057/en/IDC-FutureScape-2026-Predictions-Reveal-the-Rise-of-Agentic-AI-and-a-Turning-Point-in-Enterprise-Transformation>
4. Bank of England & Financial Conduct Authority. (2024). Artificial intelligence in UK financial services — 2024.
Bank of England, November 2024
<https://www.bankofengland.co.uk/report/2024/artificial-intelligence-in-uk-financial-services-2024>
5. McKinsey & Company. (2025). The state of AI in 2025: Agents, innovation, and transformation.
QuantumBlack, AI by McKinsey, November 2025.
<https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai>
6. McKinsey & Company. (2025). Seizing the agentic AI advantage.
QuantumBlack, AI by McKinsey, June 2025.
<https://www.mckinsey.com/capabilities/quantumblack/our-insights/seizing-the-agentic-ai-advantage>

AUTHOR & CONTRIBUTOR

Andrea Perl

Partner & Principal Data Consultant,
Zühlke

PRODUCTION

Sachin Kharchane

Graphic Designer,
GFTN

