

PAARA

Privacy-Aware AI Risk Architecture Connector-DPIA Methodology v1.3

*A six-component methodology for assessing privacy and security risks
at the AI connector layer in enterprise GenAI deployments*

Vivek Kumar

Responsible AI and Data Privacy Governance | Indiana Wesleyan University
AIGP | CIPP/US | CIPP/E | CIPM | CISM

Version 1.3 | May 2026

SSRN: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=6784078

This work is non-confidential. No employer, client, or system-specific information is disclosed.

Abstract

Enterprise generative AI tools — Microsoft 365 Copilot, Google Gemini for Workspace, ChatGPT Enterprise, and systems built on Model Context Protocol (MCP) servers — operate through connector integrations: OAuth scopes, Microsoft Graph permissions, REST API bridges, and retrieval-augmented generation pipelines that link a large language model to enterprise data sources including email, calendars, document repositories, HR systems, finance platforms, and customer databases.

Major data protection authorities — the EDPB, the UK ICO, the CNIL, NIST, Singapore IMDA, the Australian OAIC, and Canada OPC — have each published guidance on AI and privacy obligations, and confirm that AI deployments generally trigger mandatory DPIA requirements. The consistent regulatory position is that a system-level DPIA must encompass all relevant data flows, including connector-layer processing. Current guidance does not, however, provide practitioners with an operational methodology for assessing those connector-layer data flows. The integration layer — the OAuth scope, the permission, the retrieval boundary, the prompt-injection vector, the insider-accessible data surface — is acknowledged as in scope but is left without structured assessment guidance. Most system-level DPIAs conducted in practice do not address it adequately, leaving the DPIA incomplete precisely where regulators require completeness.

This paper introduces the PAARA Connector-DPIA Methodology: a six-component practitioner module designed to be embedded within a system-level DPIA to ensure the connector layer is assessed completely. The methodology does not create a separate legal instrument or duplicate existing DPIA obligations. It provides the structured assessment content — connector inventory, permission risk register, retrieval threat model, insider-risk assessment, risk scoring, and cross-framework alignment — that system-level DPIAs currently lack for the connector layer. In the ICO's own framing, a system-level DPIA must encompass all relevant data flows. PAARA is the methodology that makes that requirement achievable in practice for enterprise AI connector deployments.

The methodology maps enterprise connector deployments against the WP29 nine-criteria DPIA trigger framework and aligns findings with GDPR Article 35, the EU AI Act Article 27 FRIA obligation enforceable from August 2026, and NIST AI 600-1, as well as ISO/IEC 42005:2025, OWASP LLM Top 10, and related standards. The objective is a single connector assessment document that satisfies multiple overlapping regulatory and governance requirements without requiring separate reviews for each framework.

Keywords: DPIA; AI connector risk; GDPR Article 35; EU AI Act; ISO/IEC 42005; NIST AI RMF; Microsoft 365 Copilot; enterprise GenAI governance; OAuth scope risk; prompt injection; insider risk; privacy-aware AI architecture; retrieval-augmented generation.

1. The Governance Gap at the Connector Layer

1.1 What current frameworks address — and what they do not

In most enterprise deployments, the privacy team reviews the AI vendor. The IT team configures the connector. Nobody reviews both together. That is the gap this methodology addresses.

Enterprise AI governance normally begins with the AI system, vendor, model, or business use case. Privacy assessments ask whether processing is lawful, necessary, proportionate, transparent, and rights-respecting. Regulators are clear that AI deployments generally require a DPIA, and that system-level DPIAs must encompass all relevant data flows — including connector-layer processing. Each of these positions is correct and necessary. What is missing is the operational methodology to actually assess those connector-layer data flows within a system-level DPIA.

The gap is not that existing laws and standards ignore privacy, AI impact, or security. They clearly address those areas. The gap is that current guidance provides no structured operational content for the connector layer as an assessment object: no inventory framework, no permission risk register, no retrieval-boundary threat model, no insider-risk assessment. System-level DPIAs conducted in practice routinely omit the connector layer entirely — not because practitioners are negligent, but because no published methodology tells them how to assess it. PAARA fills that gap.

Existing assessment lens	Primary strength	Connector-layer gap PAARA addresses
GDPR DPIA	Rights-and-freedoms risk, necessity and proportionality, safeguards, accountability	DPIA templates often assess the processing activity broadly. PAARA adds concrete questions about AI retrieval paths, inherited permissions, source reconstruction, prompt-triggered access, and cross-repository synthesis.
AI impact assessment / FRIA	Affected persons, foreseeable harms, human oversight, high-risk AI context	AI assessments often focus on system function and downstream impact. PAARA asks how connectors change who and what the AI system can reach.
AI management system	Governance roles, lifecycle controls, risk processes, documented accountability	Management-system controls need implementation artefacts. PAARA creates connector inventories, risk registers, threat models, and review evidence.
Cybersecurity review	Access control, logging, monitoring, vulnerability management, incident response	Security reviews may validate technical controls without evaluating privacy-specific aggregation, inference, data-subject rights, and insider-risk amplification.
Vendor due diligence	Contractual, processing, security, and sub-processor review	Vendor review may not capture tenant-specific permission hygiene, repository scope, user behaviour, or local data-governance maturity.

1.2 Why this gap is actionable now

Three developments converge in 2026. First, enforcement: the Italian Garante fined OpenAI €15 million in December 2024, partly on DPIA grounds, confirming that regulators will use DPIA obligations as an enforcement lever for AI deployments.

Second, the EU AI Act Article 27 FRIA obligation becomes enforceable on 2 August 2026.

Third — and most directly relevant to PAARA — regulators have confirmed the obligation but not the methodology. The UK ICO's Innovation Advice Service has confirmed that AI-related processing triggers DPIA requirements in the vast majority of cases, and that a system-level DPIA must encompass all relevant data flows including connector-layer processing. That regulatory position is correct. What the ICO has not published — and what no authority has published — is an operational methodology for assessing those connector-layer data flows within a system-level DPIA. PAARA provides that methodology.

1.3 Scope of this methodology

This methodology applies to the following deployment types:

- Productivity AI tools with enterprise data connectors — Microsoft 365 Copilot, Google Gemini for Workspace, ChatGPT Enterprise with file access
- Function-specific AI connectors — HR AI tools, legal AI, finance AI, code assistant AI connected to enterprise repositories
- Customer-facing AI agents with CRM or contact-centre data access
- Agentic AI systems built on MCP server architecture with one or more enterprise data source integrations
- RAG pipelines where enterprise documents are ingested into a retrieval layer accessible by a hosted or third-party LLM

The methodology does not cover standalone AI tools with no enterprise data integration, AI model development (addressed by CNIL and OAIC guidance), or AI systems with no personal data in scope.

2. PAARA Connector-Layer Risk Model

Figure 1 illustrates the structural position of the AI connector layer between enterprise data sources and the AI system. PAARA treats this connector layer as the primary assessment object — not the LLM, not the enterprise data sources, and not the user interface. The five numbered risk surfaces correspond to the five operational components of the methodology (Components B through F). Component A — the connector inventory — is the prerequisite that makes assessment of the other five surfaces possible.

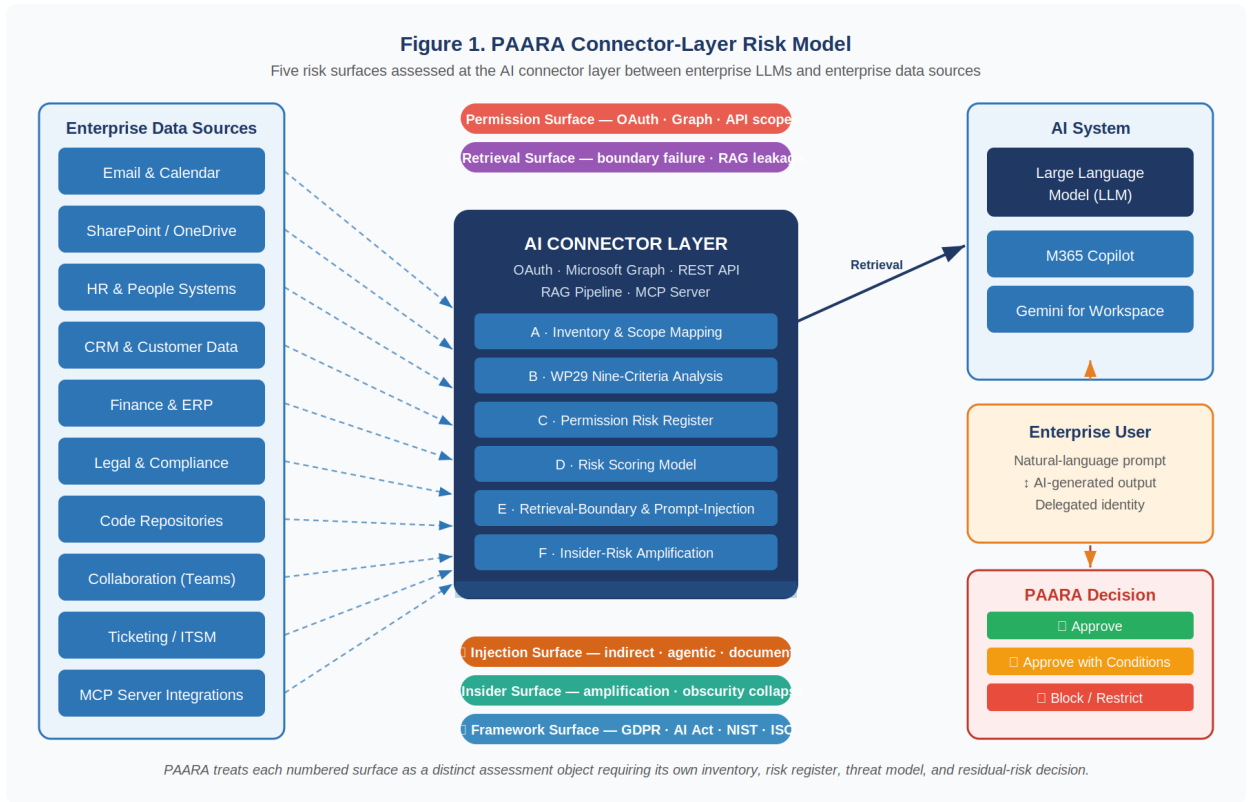


Figure 1. PAARA Connector-Layer Risk Model

The five numbered risk surfaces are assessed as distinct objects: permission scope, retrieval boundary, prompt-injection, insider amplification, and cross-framework alignment.

3. Component A — Connector Inventory and Scope Mapping

3.1 Purpose

Before any risk assessment can be conducted, the organisation must establish a complete, accurate inventory of every AI connector in scope: what it connects, what it accesses, and how access is granted. In most enterprise deployments, this inventory does not exist at the connector level. IT teams know what tools have been procured. They may not know every OAuth scope that has been granted, every Graph permission that has been configured, or every data source a connector can reach at query time.

Component A is a mandatory prerequisite for all subsequent components. An incomplete connector inventory produces an incomplete DPIA.

3.2 Connector inventory fields

Field	Description	Example
Connector ID	Unique identifier for this connector instance	CONN-001
AI Tool	Name and version of the AI product	Microsoft 365 Copilot (2025 release)
Connector type	OAuth / Microsoft Graph / REST API / RAG pipeline / MCP server / native integration	Microsoft Graph API (delegated + application permissions)
Data sources accessed	All enterprise data repositories the connector can reach	Exchange Online, SharePoint, OneDrive, Teams messages
Permission scopes granted	Full list of OAuth scopes or Graph permissions in effect	Mail.Read, Files.Read.All, Calendars.Read, Sites.Read.All
Personal data categories	Types of personal data accessible through the connector	Employee communications, HR documents, client correspondence
Special category data?	Does the connector reach Art. 9 special categories (health, union membership, etc.)?	Yes — employee health records in SharePoint HR site
User population	Number and categories of individuals with access	3,500 licensed users across 12 business units
Third-party processor?	Does the connector route data to a third-party service outside the tenant?	Yes — OpenAI API (DPA in place)
Cross-border transfers?	Does data flow across jurisdictional boundaries?	Yes — EU data processed in US (SCCs in place)
DPIA owner	Named accountable owner for this connector's DPIA	DPO / CISO joint ownership

3.3 Scope-minimisation assessment

For each connector, assess whether the permission scopes granted are the minimum necessary for the declared use case.

The most common over-permissioning pattern is simply that the broadest available scope was selected during initial setup because it was the easiest option. Files.Read.All is one click. Scoping to a specific site collection requires configuration. Mail.Read at tenant level is the default. The result is that most enterprise AI connectors are running on significantly broader permissions than their declared use case requires.

Specific patterns to check:

- Files.Read.All granted when only a specific SharePoint site is needed
- Mail.Read granted at tenant level when only a single mailbox is required
- MCP server granted write permissions where the use case requires read only

4. Component B — WP29 Nine-Criteria Analysis at the Connector Layer

4.1 Framework

The WP29 Guidelines on DPIA (wp248rev.01, October 2017, endorsed by EDPB May 2018) provide nine criteria for determining when a DPIA is mandatory. Where processing meets two or more criteria, a DPIA is generally required.

#	WP29 Criterion	Application to AI connector layer	Typical connector finding	Triggered?
1	Evaluation or scoring	Connectors enabling AI to evaluate employee performance, productivity, or behaviour from communications and calendar data trigger this criterion.	HR AI connectors; productivity analytics	Conditional
2	Automated decision-making with legal/significant effect	Where connector outputs feed directly into access decisions, credit decisions, or HR actions without human review, Article 22 is triggered alongside the DPIA obligation.	Loan-decisioning AI; automated HR workflow connectors	Conditional
3	Systematic monitoring	An AI connector that continuously indexes and retrieves employee communications constitutes systematic monitoring of workplace behaviour — even where individual queries are user-initiated.	Copilot/Gemini with Mail.Read at tenant level	YES

4	Sensitive data or highly personal nature	Where the connector provides access to SharePoint HR sites, health data, union correspondence, or financial records, criterion 4 is met regardless of whether such data is the primary use case.	Most enterprise tenants — sensitive data co-located with general content	YES (most)
5	Data processed on a large scale	An enterprise-wide connector licence constitutes large-scale processing. The WP29 volume, variety, and velocity test is met by the connector's breadth of access, not merely by the number of queries.	Any enterprise-wide deployment	YES
6	Matching or combining datasets	This criterion is triggered structurally by AI connector architecture: a connector combines previously siloed data — email, calendar, SharePoint, HR, CRM — at query time. This combination did not exist before the connector was deployed.	ALL productivity AI connectors by design	YES — structural
7	Data concerning vulnerable subjects	Where the connector accesses patient data, student records, or beneficiary information, criterion 7 is met.	Healthcare, education, public sector deployments	Conditional
8	Innovative use or new technologies	The ICO states that AI-related processing is generally a high-risk activity that requires a DPIA. Current enterprise AI connector deployments are novel relative to prior AI use cases.	All enterprise GenAI connector deployments	YES
9	Prevents data subjects exercising a right or using a service	Where a connector-enabled AI makes recommendations that affect access to services or employment without adequate human review, criterion 9 is met.	Access control AI; workforce management connectors	Conditional

4.2 DPIA obligation determination

Most enterprise AI connector deployments of general-purpose productivity tools will satisfy criteria 3, 4, 5, 6, and 8 simultaneously, making the DPIA obligation clear under GDPR Article 35 and UK GDPR equivalent provisions. Practitioners should document the criteria analysis for each connector and record the determination explicitly. Where an organisation believes a specific connector does not trigger the DPIA threshold, that reasoning should be written down and retained — it will be the first thing a regulator asks for.

5. Component C — OAuth/Graph Permission Risk Register

5.1 Purpose and framing

The permission configuration of an AI connector is the primary determinant of its data exposure profile. An AI tool granted `Files.Read.All` can retrieve any document in the tenant. The same tool scoped to a single SharePoint site collection presents a fundamentally different risk profile. The difference is not in the AI model — it is entirely in the connector scope.

Five principles govern the permission review. The first two are the most commonly violated in practice: least privilege (access only what the use case actually requires) and purpose limitation (each scope must map to a documented, specific business purpose, not a general one). The others follow from these — preserve user-level access boundaries wherever the architecture allows, separate retrieval permissions from action permissions so that read access does not silently enable write, send, or delete operations, and ensure all permissions are revocable, auditable, and subject to periodic recertification.

In practice, the second and third principles are where most deployments fail. Broad scopes are granted at deployment because narrower scopes require additional configuration work. The recertification cycle is rarely established at the time of deployment and is never retroactively applied. This means that the permission profile of a connector deployed twelve months ago may bear no relationship to the use case it was deployed to support.

High-risk permission patterns to flag immediately:

- Application-level permissions (rather than delegated/user-context permissions) — these allow the connector to access data the querying user cannot directly access
- Write, send, delete, or workflow-execution permissions granted alongside retrieval permissions without separate approval
- Any permission scope that covers the entire tenant rather than a named group, site, or mailbox

Permission scope	Data accessible	Minimisation compliant?	Special category risk?	Risk level	Recommended control
Mail.Read (tenant-wide)	All email content and metadata for all users	NO — typically over-broad	Yes — likely	HIGH	Scope to named mailboxes; apply sensitivity label exclusions; DLP on AI outputs
Files.Read.All	All files across all SharePoint sites and OneDrive	NO — typically over-broad	Yes — likely	HIGH	Scope to specific site collections; Restricted SharePoint Search; sensitivity label blocking

Calendars.Read	Calendar entries for all users	Conditional	Low — unless medical appointments visible	MEDIUM	Acceptable for productivity use cases with user-consent notice
Sites.Read.All	All SharePoint site content	NO — typically over-broad	Yes — likely	HIGH	Apply Restricted SharePoint Search; implement sensitivity label policy; quarterly scope review
User.Read.All	Full user profile directory	Conditional	Low for standard directory data	LOW-MEDIUM	Generally acceptable; review if profile includes disability or health data
MCP server with write permissions	Ability to create, modify, or delete enterprise data	HIGH RISK — requires explicit justification	Yes — data integrity risk	CRITICAL	Human-in-the-loop approval for all write actions; log every write; implement rollback

5.2 Sensitivity label integration

Where the enterprise operates a Microsoft Purview or equivalent sensitivity labelling regime, the connector DPIA must document how sensitivity labels interact with the connector's access scope — specifically which labels block AI access, which do not, and what percentage of sensitive content carries a label. Unlabelled content is the primary vector through which special-category data reaches AI connectors in enterprises with partial labelling coverage.

6. Component D — Connector Risk Scoring Model

6.1 Purpose

The scoring model converts the findings from Components A through C into a deployment decision. It is a governance tool, not a compliance test. A connector that scores LOW still needs a DPIA if the nine-criteria analysis in Component B requires one. A connector that scores CRITICAL should not be deployed regardless of what business pressure exists. The scoring guides the conversation; it does not replace legal or privacy professional judgment.

6.2 Risk factor table

Risk factor	Weight	Scoring prompt
Sensitive personal data access	High	Can the connector retrieve special-category, HR, health, legal, financial, or protected-class data?

Cross-system aggregation	High	Can the AI synthesise records across repositories that were historically separate?
Application or service-account permissions	High	Can the connector retrieve records outside individual user-level access?
HR/legal/security repository access	High	Are repositories with employment, privilege, investigation, or disciplinary records reachable?
Lack of retrieval-level logging	High	Can the organisation reconstruct what source records were retrieved for each output?
Agentic action capability	Critical	Can the AI send, update, approve, delete, route, or trigger workflows?
External sharing exposure	Medium	Can outputs or retrieved content flow to guests, partners, or outbound communications?
User population size	Medium	How many users can access the connector and how diverse are their roles?
Prompt-injection exposure	High	Can retrieved content contain instructions that affect AI behaviour?

6.3 Decision rules

Score outcome	Decision rule	Recommended decision
LOW	No critical factors; no more than one medium factor; logging and least privilege confirmed	Approve with standard monitoring
MODERATE	One high factor or multiple medium factors; controls available and owner accepts residual risk	Approve with conditions
HIGH	Two or more high factors, sensitive repositories, cross-system aggregation, or weak logging	Restrict or defer pending full connector-DPIA
CRITICAL	Any agentic action capability combined with sensitive data, broad permissions, or missing logging	Block or suspend until remediated

7. Component E — Retrieval-Boundary and Prompt-Injection Threat Model

7.1 Retrieval-boundary risk

A retrieval boundary is the set of documents, records, or data objects that an AI connector can retrieve and include in its context window in response to a user query. In a misconfigured or over-permissioned deployment, the retrieval boundary is determined by the connector's OAuth scope, not by the user's individual access rights. This creates a class of privacy failure that is architecturally distinct from conventional access control failure: the connector can retrieve and synthesise data the querying user does not have direct permission to access, because the connector itself holds the access right.

Retrieval-boundary risk	Mechanism	Control
Scope exceeds user permissions	Connector application permission grants the AI access to files the querying user cannot directly open	Use delegated permissions (user context) not application permissions where possible
Cross-department data leakage	Query about Project X retrieves documents from HR, Legal, or Finance sites that reference the project but are outside the user's role scope	Implement site-level exclusions; apply Restricted SharePoint Search; conduct retrieval-boundary testing before deployment
RAG pipeline boundary failure	Ingested document corpus does not enforce access controls at retrieval time — all ingested content is available to all queries regardless of user identity	Require per-user retrieval filtering; retest after every corpus update; audit retrieval logs for anomalous cross-boundary results
Agentic connector over-reach	Agentic AI accumulates context from multiple data sources in a single session beyond what any single query would retrieve	Implement session-scope limits; require explicit user confirmation for cross-connector data aggregation

7.2 Prompt-injection threat model

Prompt injection is the injection of attacker-controlled instructions into the AI's context window via content that the AI retrieves or processes. NIST AI 600-1 is the only major framework that explicitly acknowledges prompt injection can exploit vulnerabilities in interconnected systems.

- Direct injection: user submits a malicious prompt designed to extract data beyond their authorised access
- Indirect injection via ingested content: a document accessible to the connector contains instructions designed to manipulate the AI's response when that document is retrieved
- Indirect injection via email: an externally sent email contains instructions causing a mail-connected AI to forward sensitive content to an external address

- Agentic multi-step injection: injection in early-stage retrieved content redirects the AI's subsequent action steps

8. Component F — Insider-Risk Amplification Assessment

8.1 The insider-risk amplification effect

AI connectors reduce the friction of accessing, synthesising, and exfiltrating large volumes of enterprise data for a user who has legitimate connector access. Before an AI connector existed, an employee who wanted to compile a comprehensive view of customer contacts, financial projections, or HR records would need to navigate multiple systems, export files individually, and aggregate content manually. That process took hours. It left traces. It was noticeable.

An AI connector changes that. A user with a legitimate Copilot licence can, in a single session, query for customer contact lists, financial summaries, HR headcount data, and competitive intelligence — and receive synthesised, structured output. The connector has not created a new access right. It has made the exercise of an existing access right fast, clean, and hard to distinguish from normal work.

This is the single risk dimension that no data protection authority surveyed has specifically addressed. Every other risk in this methodology has at least a partial analogue in existing guidance. Insider-risk amplification does not. It requires a different framing: the threat is not that the connector was misused — it is that the connector was used exactly as designed.

Assessment dimension	Questions to address	Indicative controls
Data aggregation potential	What is the maximum personal data volume a single licensed user could compile in one session via normal queries?	Session-level data retrieval volume limits; anomaly detection on bulk-compilation query patterns
Exfiltration surface	Can AI outputs be exported, copied, forwarded externally, or shared outside the tenant without DLP controls?	DLP policy on AI-generated outputs; block external forwarding of AI summaries; sensitivity label inheritance
Privilege-to-connector alignment	Does the user's connector access correspond to their actual job role, or does the connector grant access beyond functional need?	Role-based connector access tiers; quarterly access review; just-in-time elevation for high-sensitivity connector access
Audit trail completeness	Are all connector queries, retrieved documents, and AI-generated outputs logged? Is the audit trail tamper-resistant?	Mandatory audit logging; SIEM integration; UEBA on connector query patterns
Offboarding and access revocation	Are connector access rights immediately revoked on employee termination or role change?	Connector access tied to active employment status in IAM; pre-offboarding AI activity review for high-risk roles

High-risk role identification	Which roles combine the broadest connector access with the highest potential exfiltration motivation?	Enhanced monitoring for high-risk roles; reduced connector scope for employees under HR investigation
--------------------------------------	---	---

9. Component G — Cross-Framework Alignment

9.1 Purpose

Privacy teams in global enterprises are not running one assessment — they are running several simultaneously, often without realising it. A Copilot deployment in an EU entity requires a GDPR DPIA. The same deployment in a high-risk AI classification may require an EU AI Act FRIA from August 2026. The CISO wants alignment to NIST AI 600-1. The ISO 42001 certification programme wants evidence of an impact assessment. Component G maps the connector findings from the earlier components to each of these frameworks so that the work is done once, not four times.

Connector risk	GDPR Art. 35	EU AI Act Art. 27	ISO/IEC 42005:2025	NIST AI 600-1	OWASP LLM Top 10
OAuth/Graph over-permissioning	Art. 5(1)(c) data minimisation; DPIA item	Relevant where HR connectors trigger Annex III	Clause 6.1.2 — technical safeguards	GOVERN 1.7; MANAGE 2.2	LLM07: Insecure Plugin Design
Retrieval-boundary failure	Art. 5(1)(f) integrity and confidentiality	Fundamental rights to data protection	Clause 8.4 — technical controls	INFO SECURITY — retrieval integrity	LLM03: Excessive Agency; LLM06: Sensitive Information Disclosure
Prompt injection via ingested content	Art. 32 — security of processing	Integrity of AI system output	Annex B — adversarial inputs	EXPLICIT — prompt injection in interconnected systems	LLM01: Prompt Injection
Insider-risk amplification	Art. 32(4) — technical and organisational measures	Not directly in FRIA scope	Clause 6.1.2 — misuse scenarios	Not explicitly addressed — PAARA gap contribution	LLM08: Excessive Permissions
Dataset combination at query time	WP29 criterion 6 — triggers mandatory DPIA	Relevant where combined profile used for FRIA-covered decisions	Clause 6.1.3 — scope of impact assessment	MAP 1.1 — context-specific risks	LLM02: Sensitive Information Disclosure
Absence of DPIA	Art. 83(4) — up to €10M or 2% annual	Art. 27 non-compliance —	ISO 42001 Clause 6.1.4 —	GOVERN documentation	Governance — no direct Top 10 mapping

	global turnover	enforceable 2 August 2026	mandatory for certified orgs	on requirement	
--	-----------------	---------------------------	------------------------------	----------------	--

9.2 FRIA applicability determination

The EU AI Act Article 27 FRIA obligation becomes enforceable on 2 August 2026. It is not co-extensive with the GDPR DPIA obligation. Practitioners should make an explicit FRIA applicability determination for each connector: (1) does the AI system fall within Annex III high-risk categories, particularly §4 (employment and workers management), §5 (access to essential private services), or §2 (critical infrastructure)?; (2) is the deployer a public body or private entity providing a service of a public nature? Document the determination with reasoning. Where FRIA applies, prepare a consolidated DPIA/FRIA document. Where it does not, document the basis — it will be needed if the question is ever raised.

10. Worked Example — Enterprise AI Assistant Connected to Collaboration and HR Systems

This example is entirely fictional and non-confidential. It is designed to illustrate how the methodology is applied in practice. No employer, client, or system-specific information is reflected.

10.1 Scenario

A mid-sized professional services organisation is piloting an AI assistant connected to SharePoint, Outlook, Teams, and selected HR knowledge-base content for a workforce productivity pilot of 800 licensed users.

Assessment step	Finding	PAARA decision impact
A. Scope	AI assistant connected to SharePoint, Outlook, Teams, and selected HR knowledge-base content for an 800-user productivity pilot.	Connector is not merely technical. It changes retrieval across collaboration and HR-adjacent records.
A. Data categories	Business documents, email content, chat content, calendar metadata, employee policy documents, and some employee-relations references are reachable.	Personal data and employment-context data are present. DPIA threshold screening required.
A. Identity model	System uses user-delegated access for SharePoint and Outlook, but a service account indexes certain knowledge-base materials.	Delegated access lowers risk. Service-account indexing requires a separate permission review.
B. WP29 criteria	Systematic monitoring (criterion 3), combining datasets (criterion 6), innovative technology (criterion 8),	Five criteria triggered. Full connector-DPIA addendum required.

	employment-context (criterion 4 and 7), and possible evaluation/scoring (criterion 1) are all implicated.	
C. Permission risk	Read scopes are justified for approved repositories, but chat and email retrieval increase sensitivity. HR knowledge-base access is overbroad for the pilot scope.	Restrict HR knowledge-base access; approve collaboration access only for the defined pilot group.
D. Risk score	Two high factors (HR data reachable; cross-system aggregation via Teams+SharePoint+Outlook). Logging not yet enabled. No agentic action capability.	Score: HIGH. Defer general rollout pending remediation. Pilot may proceed with conditions.
E. Retrieval threat model	Prompt tests show queries seeking employee complaints, manager comments, and restructuring references return results from HR knowledge-base documents.	Block sensitive employee-targeted prompts. Exclude HR investigation and legal folders.
F. Insider-risk amplification	Natural-language querying collapses practical obscurity. A single user can synthesise across messages, documents, and calendars in one session.	Enhanced monitoring and acceptable-use notice required before pilot users are activated.
Deployment decision	Pilot may proceed only after HR/legal repositories are excluded, retrieval logs are enabled, prompt-abuse monitoring is configured, and users receive notice.	APPROVE WITH CONDITIONS. Review after 60 days.

This outcome was not straightforward to reach. The business owner initially challenged the HR knowledge-base restriction, arguing that employee policy documents were central to the productivity use case and that removing them would significantly limit the tool's value. The assessment team held the position: the HR repository contained employee-relations references that were out of scope for a productivity pilot and could not be adequately segregated at the time of deployment. The resolution was a 30-day deferral of that data source, with a commitment to re-scope and retest before it was re-included. That kind of negotiation is normal. The methodology provides the framework for having it on documented grounds rather than as an informal disagreement between IT and legal.

Deployment decision text: The connector may be used for a limited productivity pilot. It may not retrieve HR investigations, legal files, privileged advice, security vulnerability records, disciplinary material, employee complaints, health records, or compensation data. The organisation must enable retrieval-level logging, apply sensitive-prompt blocking, monitor repeated person-focused queries, and recertify permissions before any expansion.

11. Limitations

This paper is a practitioner methodology, not legal advice. It does not replace a formal DPIA, FRIA, security assessment, vendor review, or legal analysis where those are required.

The methodology is designed for enterprise GenAI connector governance and may require sector-specific adaptation for healthcare, financial services, education, public sector, or law-enforcement contexts. The risk scoring model is a governance tool — scores guide deployment decisions but do not constitute a legal compliance determination.

The cross-framework alignment table reflects the state of published guidance as of May 2026 and will require updating as ICO, CNIL, NIST, and IMDA publish forthcoming connector-specific guidance. Practitioners should treat this as a living document and revalidate the framework alignment annually or when material regulatory developments occur.

Nothing in this methodology discloses employer-confidential information. All examples are fictional and illustrative. The views expressed are the author's own.

12. Future Research Directions

- Connector certification profiles for approved, restricted, and prohibited enterprise AI connectors
- Agentic AI governance where retrieval is coupled with action, workflow execution, messaging, or transaction authority
- Retrieval-aware access controls that distinguish technical access from policy-permitted AI use
- AI-specific insider-risk metrics for natural-language discovery, cross-context aggregation, and inference acceleration
- Sector-specific PAARA profiles for HR, legal, healthcare, financial services, education, public sector, and software engineering
- Regulatory harmonisation between DPIAs, FRIAs, AI impact assessments, cybersecurity controls, and internal audit assurance
- Empirical benchmarking of connector risk scores against observed security incidents and regulatory enforcement actions

Conclusion

Enterprise AI governance cannot focus only on models, vendors, or high-level use cases. As organisations move from experimentation to scaled deployment, connectors become one of the most important control points in the AI operating model. They determine what the AI system can see, retrieve, combine, infer, disclose, and act upon.

The PAARA Connector-DPIA Methodology provides the operational content that system-level DPIAs currently lack for the connector layer. It does not create a new legal obligation or a separate legal instrument. Regulators have already established the obligation — a system-level DPIA must encompass all relevant data flows. PAARA makes that obligation achievable in practice by providing the structured methodology to assess connector-layer data flows completely: inventory, permission risk, retrieval boundaries, prompt-injection threats, insider-risk amplification, and cross-framework alignment.

The central governance question regulators have already answered: if an AI system processes personal data through a connector, that processing must be captured in the DPIA. What regulators have not yet provided is an operational methodology for doing so. Most organisations deploying Microsoft 365 Copilot, Google Gemini for Workspace, or agentic AI on MCP server architectures today cannot demonstrate that their system-level DPIA adequately addresses what the connector can reach, why it can reach it, how retrieval is constrained, how misuse is detected, and who accepted the residual risk. PAARA exists to close that gap — not by creating a new legal obligation, but by making an existing one achievable.

Appendix A — Practitioner Checklist

Answer each question before approving deployment. Any 'No' on questions 1–10 should defer or restrict production deployment until resolved.

#	Question	Answer (Yes / No / N/A)
1	Do we know every data source the connector can reach?	
2	Do we know whether personal data or sensitive data is reachable?	
3	Do we understand the identity and permission model?	
4	Have we reviewed OAuth, Graph, API, or service-account scopes for least privilege?	
5	Have we tested whether user-level permissions are accurately enforced at retrieval time?	
6	Have we assessed whether the connector triggers DPIA high-risk criteria (WP29 nine-criteria)?	
7	Have we defined and tested the retrieval boundary?	
8	Have we tested prompt-injection and malicious-content scenarios?	
9	Have we assessed insider-risk amplification?	
10	Can we reconstruct which source records were retrieved for a given output?	
11	Can users or data subjects receive meaningful transparency where required?	
12	Can we detect suspicious AI-mediated access patterns?	
13	Have we restricted high-risk repositories (HR investigations, legal files, security vulnerability records)?	
14	Have we documented residual risk and obtained named owner acceptance?	
15	Have we set a review date and material-change reassessment triggers?	

Appendix B — One-Page Connector-DPIA Intake Form

Field	Response
Connector name	
AI system	
Business purpose	
Business owner	
Technical owner	
Privacy / DPO owner	
Security owner	
Connected data sources	
Data categories	
Sensitive data present?	
User groups	
Identity model (delegated / service account / application permission)	
Permission scopes	
Retrieval model	
Output model	
Logging available?	
Source reconstruction possible?	
Prompt-injection testing completed?	
Insider-risk review completed?	
WP29 criteria triggered	
DPIA threshold result	
Required controls	
Residual risk rating	
Deployment decision	
Review date / reassessment triggers	

References

[1]	Regulation (EU) 2016/679 (GDPR), Article 35 — Data Protection Impact Assessment.
[2]	Article 29 Working Party, Guidelines on DPIA and determining whether processing is 'likely to result in a high risk', WP248 rev.01, October 2017, endorsed EDPB May 2018.
[3]	EDPB Opinion 28/2024 on certain data protection aspects related to AI models, 17 December 2024.
[4]	EDPB ChatGPT Taskforce Report, 23 May 2024.
[5]	UK ICO Generative AI Consultation Series (six chapters), 2024.
[6]	UK ICO Tech Futures: Agentic AI, 2026.
[7]	CNIL Recommendations on the Development of AI Systems, April 2024.
[8]	CNIL AI How-To Sheets, January 2026.
[9]	Regulation (EU) 2024/1689 (EU AI Act), in force 1 August 2024; Article 27 FRIA enforceable 2 August 2026.
[10]	NIST AI Risk Management Framework 1.0 (AI RMF), NIST AI 100-1, January 2023.
[11]	NIST Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile, NIST AI 600-1, 26 July 2024.
[12]	NIST Security and Privacy Controls for Information Systems and Organizations, SP 800-53 Rev. 5.
[13]	NIST AI RMF Critical Infrastructure Profile Concept Note, 7 April 2026.
[14]	PDPC/IMDA Model AI Governance Framework for Agentic AI, January 2026.
[15]	OAIC Guidance on Privacy and Commercially Available AI Products, October 2024.
[16]	ISO/IEC 42005:2025 — AI System Impact Assessment.
[17]	ISO/IEC 42001:2023 — AI Management System.
[18]	OWASP Top 10 for Large Language Model Applications 2025, v2.0, OWASP GenAI Security Project, 17 November 2024.
[19]	OWASP AI Security and Privacy Guide (AI Exchange), owaspai.org, 2022–2026.
[20]	Italian Garante v. OpenAI, Provvedimento del 20 dicembre 2024, €15M administrative fine.
[21]	Microsoft Build Your Own DPIA Templates for Microsoft 365 Copilot, Microsoft Corporation, 2024.
[22]	DSK Statements on Microsoft 365 Copilot DPIA Obligations, Datenschutzkonferenz (DSK), 2024–2025.

Endnotes

- [1] Italian Garante v. OpenAI, Provvedimento del 20 dicembre 2024, €15M administrative fine for GDPR violations including inadequate DPIA.
- [2] EU AI Act (Regulation (EU) 2024/1689), Article 27 — Fundamental Rights Impact Assessment for High-Risk AI Systems, enforceable 2 August 2026.
- [3] GDPR Article 5(1)(c): personal data shall be 'adequate, relevant and limited to what is necessary in relation to the purposes for which they are processed' (data minimisation).
- [4] Article 29 Working Party, Guidelines on Data Protection Impact Assessment (DPIA) and determining whether processing is 'likely to result in a high risk' for the purposes of Regulation 2016/679, WP248 rev.01, 4 October 2017, endorsed by EDPB 25 May 2018.
- [5] NIST AI 600-1, Generative Artificial Intelligence Profile, 26 July 2024, Information Security risk area: 'Prompt injection can use generative AI to exploit vulnerabilities in interconnected systems.'

About the Author

Vivek Kumar is a Responsible AI and Data Protection governance professional with expertise in privacy-aware AI risk architecture, enterprise GenAI controls, AI connector risk governance, and regulatory-aligned AI operating models. He holds the AI Governance Professional (AIGP), CIPP/US, CIPP/E, CIPM, and CISM certifications and is affiliated with Indiana Wesleyan University. His published work includes peer-reviewed contributions on AI governance and privacy risk architecture, with work available on SSRN and ResearchGate.

This methodology is derived from practitioner experience and publicly available regulatory analysis. It does not disclose confidential employer information. The views expressed are the author's own.

Citation: Kumar, Vivek. PAARA Connector-DPIA Methodology v1.3: A Privacy-Aware Risk Assessment Model for Enterprise GenAI Connectors. Practitioner methodology, May 2026. SSRN: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=6784078

Appendix C — Version History

Version	Date	Changes
1.0	20 May 2026	Initial six-component practitioner methodology draft
1.1	20 May 2026	Added governance-gap framing table, risk scoring model, worked example, limitations, future research directions, and reference list
1.2	21 May 2026	Added Figure 1 architecture diagram; added Component D risk scoring model; added worked example (Section 10); added limitations; added future research; added practitioner checklist (Appendix A); added intake form (Appendix B); added endnotes; revised abstract framework list; humanised prose — added practitioner observations, friction sentence in worked example, replaced hedging language with direct statements, converted five-bullet principles list to prose
1.3	21 May 2026	Revised abstract, Section 1.1, Section 1.2, and conclusion to align with UK ICO Innovation Advice Service position (May 2026). PAARA repositioned as a connector-layer assessment module embedded within system-level DPIAs, not a separate legal instrument. ICO confirmed that system-level DPIAs must encompass all connector-layer data flows — PAARA provides the operational methodology to achieve that requirement in practice.