

2026 년 8 월 24 일

「인공지능 안전성 확보 의무 이행방법에 관한 고시」 제정안에 관한 BUSINESS SOFTWARE ALLIANCE(BSA) 의견서

과학기술정보통신부 귀하

BSA 비즈니스 소프트웨어연합(The Business Software Alliance, **BSA**)¹은 「인공지능 발전과 신뢰 기반 조성 등에 관한 기본법」(이하 '인공지능기본법') 제32조²를 이행하는 「인공지능 안전성 확보 의무 이행방법에 관한 고시」 제정안(이하 '고시(안)')과 관련하여, 과학기술정보통신부에 다음과 같이 의견을 제출합니다.

BSA는 각국 정부와 세계 시장을 중심으로 글로벌 소프트웨어 산업을 대변하고 있는 연합체로, 회원사들은 AI 제품 및 서비스를 제공할 뿐만 아니라, 제3자가 AI 시스템 및 애플리케이션 개발에 활용할 수 있는 도구를 제공하는 데 앞장서고 있습니다.

BSA는 인공지능기본법 및 그 하위법령의 마련 과정 전반에 걸쳐 과학기술정보통신부와 지속적으로 소통해 왔습니다. 아울러 BSA는 인공지능기본법이 한국의 AI 정책 목표를 보다 충실히 뒷받침할 수 있도록 이를 정비하기 위한 과학기술정보통신부의 지속적인 노력에도 협력해 왔습니다³. 이처럼 인공지능기본법을 정비하기 위한 노력이 진행 중임을 고려하여, 해당 작업이 마무리될 때까지 고시(안)의 제정을 유예해 주실 것을 요청드립니다. 이와 함께 다음의 내용으로 고시 제정과 관련한 내용을 다음과 같이 제언드립니다.

요약

1. 인공지능기본법에 대한 제도 개선 논의가 진행 중인 만큼, 그 검토가 완료될 때까지 고시(안) 제정을 유예해 주시길 제언드립니다. 현시점에 구속력 있는 규정을 제정하는 것은 인공지능기본법 제도개선 연구반의 소관에 명백히 속하는 쟁점을 선점하게 되고, 이미 복잡한 규제 환경에 또 하나의 규범을 추가하게 됩니다. 또한 빠른 시일 내 개정할 수 있는 의무 사항을 기업이 미리 대비하도록 요구하는 결과를 초래할 수 있다는 점이 우려됩니다.

¹BSA 회원사: Adobe, Alteryx, Amadeus, Amazon Web Services, Asana, Atlassian, Autodesk, Avalara, Bentley Systems, Box, Cisco, Cloudflare, Cohere, Cohesity, Dassault Systemes, Databricks, Datadog, Docusign, Dropbox, Elastic, EY, Graphisoft, HubSpot, IBM, Kyndryl, MathWorks, Microsoft, Notion, Okta, OpenAI, Oracle, PagerDuty, Palo Alto Networks, PTC, Rubrik, Salesforce, SAP, ServiceNow, Shopify Inc., Siemens Industry Software Inc., Trend Micro, TriNet, Veeam, Workday, Zendesk 및 Zoom Communications Inc.

²제 32 조(인공지능 안전성 확보 의무)는 인공지능기본법 시행령이 정한 다음의 요건이 충족되는 경우, 인공지능사업자로 하여금 AI 수명주기 전반에 걸쳐 위험을 식별·평가·완화하고, AI 관련 안전사고를 모니터링·대응하기 위한 위험관리체계를 구축하도록 요구합니다. (1) AI 학습에 사용된 누적 연산량이 10^{26} FLOPs 기준 이상일 것, (2) 해당 AI가 최첨단 기술을 적용하여 구축·운영될 것, (3) 해당 AI가 사람의 생명, 신체의 안전 및 기본권에 광범위하고 중대한 영향을 초래할 위험이 있을 것. 특히 제 32 조제 3 항은 그 구체적인 이행방법을 고시에 위임하고 있으며, 본 고시(안)이 그 대상에 해당합니다.

³참조: 「인공지능기본법 시행령 수정안에 대한 BSA 권고의견」, 2025 년 12 월, 「한국 인공지능기본법 하위법령 및 가이드라인 초안에 대한 BSA 권고의견」, 2025 년 10 월, 「인공지능기본법에 대한 BSA 의견서」, 2025 년 3 월)

2. **적용 범위를 보다 구체화해 주시길 제언드립니다.** "실질적 변경"과 "실질적 기여"를 정의하고, 제32조 의무가 AI 시스템이 아닌 AI 모델에 적용됨을 명확히 하고, AI 생태계에서 각 주체가 수행하는 역할에 따라 의무를 배분해 주시길 제언드립니다.
3. **사고 보고의 기초가 되는 정의를 정교화하고 보고 기한을 재검토해 주시길 제언드립니다.** 높은 수준의 심각성 기준을 충족하는 구체적 위해를 기준으로 "위험"과 "AI 관련 안전사고"를 정의하고, 최초 사고 보고 기한을 연장하며, 보고 의무를 부담하는 주체를 명확히 해 주시길 제언드립니다. 또한 최종 보고서는 실행 가능한 한 조속히 제출하도록 하며, 사업자가 기존의 사고 대응 체계를 활용할 수 있음을 명확히 해 주시길 제언드립니다.

규제 조항 제정의 유예 필요성

고시(안)은 인공지능기본법 제32조를 이행하기 위한 구속력 있는 규정을 마련하게 됩니다. 인공지능기본법 제32조는 다음의 요건을 모두 충족하는 AI에 적용됩니다. (1) 인공지능기본법 시행령이 정한 연산량 기준(현재 10^{26} FLOPs)을 초과하여 누적 학습되었을 것, (2) "최첨단 AI 기술"을 적용하여 "구축·운영"될 것, (3) "사람의 생명, 신체의 안전 및 기본권에 광범위하고 중대한 영향"을 초래할 위험이 있을 것(이하 '**적용대상 AI**'). 이 경우 해당 인공지능사업자는 적용대상 AI의 수명주기 전반에 걸쳐 위험을 식별·평가·완화하고, AI 관련 안전사고에 대한 위험관리체계를 구축하며, 그 결과를 과학기술정보통신부장관에게 제출해야 합니다(이하 통칭 '**제32조 의무**').

전반적으로, 현시점에 인공지능기본법과 관련된 구속력 있는 규정을 추가로 시행하는 것은 다소 이른 시점으로 사료됩니다. 2026년 3월, 과학기술정보통신부는 학계, 산업계 협회, 시민사회단체를 아우르는 전문가들이 참여하는 인공지능기본법 제도개선 연구반(이하 '**연구반**')을 공식 출범시켰습니다⁴. 과학기술정보통신부는 보도자료를 통해, 인공지능기본법 및 그 하위법령에 대한 초기 협의 과정에서 "다양한 이해관계자 간의 추가적인 협의와 공감대 형성 또는 인공지능기본법 및 시행령의 추가 개정을 요하는" "폭넓은 의견"을 접수하였다고 밝힌 바 있습니다. BSA는 연구반 참여자로서, 제32조를 포함하여 "인공지능기본법에서 제도적 개선이 필요한 영역을 발굴"하기 위해 과학기술정보통신부 및 관련 이해관계자와 협력할 기회를 갖게 된 것을 뜻깊게 생각합니다.

이러한 상황에서, 인공지능기본법에 대한 연구반의 검토 절차가 진행 중인 가운데 제32조에 관한 구속력 있는 규정을 제정하는 것은 다음과 같은 바람직하지 않은 결과를 초래할 수 있다는 점에서 우려를 갖고 있습니다.

- **개정 논의가 진행되고 있는 의무 사항에 기업이 대비해야 한다는 점을 우려하고 있습니다.** 기업은 향후 대체될 수 있는 요건을 기준으로 문서, 위험평가 체계, 사고 보고 절차를 마련해야 합니다. 이후 기업은 이미 구축한 절차를 다시 조정해야 하며, 그 결과 노력의 비효율, 비용 증가 및 추가적인 규제 혼선이 초래됩니다.
- **고시(안)은 이미 복잡한 규제 환경에 규제를 추가하는 결과를 초래할 수 있습니다.** 제32조는 현재 인공지능기본법, 시행령 및 비구속적 가이드라인에 걸쳐 규율되고 있습니다. 고시(안)은 AI 안전 의무에 관한 기존 가이드라인과 어떻게 상호작용하는지를 밝히지 않고 있는 것으로 보입니다. 하위법령이 늘어날수록 컴플라이언스의 복잡성은 증가하는 한편, 법적 확실성은 저해될 수 있다는 점이 우려됩니다.
- **현재 운영중인 인공지능 기본법 제도개선 연구반에서 논의되는 쟁점을 다루고 있습니다.** 연구반은 누적 연산량 기준이 여전히 적절한 기준인지, 의무가 AI 모델에 적용되어야

⁴[과학기술정보통신부 보도자료](#)

하는지 아니면 AI 시스템에 적용되어야 하는지, 어떤 주체가 의무를 부담해야 하는지 등 제32조와 관련된 쟁점을 검토할 것으로 예상됩니다. 고시(안)은 연구반이 자체적인 결론에 도달하기 전에 이러한 쟁점에 대한 전체를 설정하는 것으로 보일 수 있습니다.

제언: 연구반이 검토 절차를 완료하고 그에 따른 인공지능기본법 개정 사항을 확정할 때까지 고시(안)의 확정을 유예하고, 이를 반영한 수정안에 관하여 다시 의견을 수렴해 주실 것을 제언드립니다.

적용 범위의 구체화

제32조 의무는 (1) 적용대상 AI의 개발자와, (2) "실질적 변경"을 통해 AI를 적용대상 AI가 되도록 하는 모든 인공지능사업자에게 적용됩니다. 앞서 언급한 바와 같이, 10^{26} FLOPs의 누적 연산량 기준을 초과하여 학습된 AI는 적용대상으로 간주됩니다.

현행 고시(안)에서는 적용 범위를 결정하는 용어가 정의되어 있지 않고, 의무의 대상이 불분명하며, 의무의 배분이 AI 개발자와 AI 운전자(또는 이용사업자)가 실제로 수행하는 역할을 반영하지 못하고 있습니다.

- **"실질적 변경", "실질적 기여" 등 핵심 용어가 정의되어 있지 않습니다.** 고시(안) 제3조제1항제2호에 따르면, AI를 상당하게 변경하여 적용대상 AI가 되도록 한 인공지능사업자는 제32조 의무 전체를 부담하게 됩니다. 고시(안)은 어떠한 변경이 상당한 것인지 명시하지 않고 있으나, 고시(안) 제3조제2항은 "데이터 전처리 및 미세조정을 포함한 모든 학습 관련 연산"이 AI의 누적 연산량 산정에 포함된다고 규정하고 있습니다. 이는 "AI 성능의 실제 향상에 실질적으로 기여하지 않은 연산은 ... 누적 연산량에서 제외할 수 있다"고 규정하는 후속 제외 조항으로 인해 더 복잡해질 수 있습니다. 무엇이 "실질적 기여"에 해당하는지에 관한 구체성이 결여되어 있다는 점에 우려를 표합니다. 미세조정, 추가 학습, 특정 애플리케이션을 위한 모델의 조정은 일상적인 활동이며, 기업은 어떠한 활동이 제32조 의무를 촉발하는 잠재적 결과를 수반하는지를 명확히 이해할 필요가 있습니다.
- **고시(안)(및 보다 넓게는 인공지능기본법)은 제32조 의무가 AI 모델에 적용되는지, AI 시스템에 적용되는지를 명확히 규정하지 않고 있습니다.** BSA는 AI 모델(광범위한 다양한 응용 분야에 활용될 수 있는 기반 기술을 창출)과 AI 시스템(하나 이상의 AI 모델을 다른 요소와 결합하여 특정 용도의 AI 애플리케이션을 구현) 간의 구분을 인식해 주실 것을 요청드립니다. 예를 들어, 하나의 거대 언어 AI 모델이 번역 애플리케이션이나 AI 기반 챗봇과 같은 수백 개의 서로 다른 AI 시스템에 통합될 수 있습니다. 즉, 모델은 기반이 되는 기술이고, 시스템은 이를 특정 용도에 활용하는 것입니다. 학습 연산량에 초점을 두고 있음을 고려할 때, 제32조 의무는 AI 시스템이 아닌 AI 모델에 적용됨을 명확히 해 주시길 제언드립니다.
- **셋째, 고시(안)은 AI 개발자와 AI 운전자(또는 이용사업자)의 구별되는 역할을 혼동하고 있습니다.**⁵ BSA는 컴플라이언스 의무가 AI 생태계에서 각 주체가 수행하는 역할에 따라 배분되어야 하며, 각 주체는 자신이 인지하고 통제할 수 있는 범위 내의 위험에 대해서만 책임을 져야 한다고 일관되게 제언해 왔습니다. 앞서 언급한 바와 같이, 실질적 변경을 통해 적용대상이 아닌 AI를 적용대상이 되도록 한 인공지능사업자는 제32조 의무 전체를

⁵인공지능 개발자는 인공지능 시스템을 설계, 코딩 또는 제작하는 주체를 말합니다. 인공지능 활용자는 내부적으로 개발되었거나 제3자가 개발한 인공지능 시스템을 사용하는 주체를 말합니다. 하나의 주체가 개발자이면서 동시에 활용자일 수 있습니다.

부담하게 됩니다. 문제는 제32조가 운전자(또는 이용사업자)가 관여하지 않은 개발 및 학습 단계를 포함하여 AI 수명주기 전반에 걸쳐 위험을 식별·평가·완화하도록 요구한다는 점입니다. 모델을 미세조정하는 운전자(또는 이용사업자)는 기반 모델의 학습 데이터, 아키텍처, 평가 결과 또는 알려진 한계에 대한 접근 권한을 취득하는 것이 아니므로, 최초 학습 과정에서 내재된 위험을 평가할 수 있는 위치에 있지 않습니다. 그 결과, AI가 어떻게 구축되었는지에 관하여 정보가 가장 적은 당사자가 이를 문서화할 부담을 전부 지게 되는 반면, 기반 모델의 개발자는 해당 모델이 단독으로는 기준에 미달한다는 이유로 제32조상 아무런 의무를 부담하지 않게 됩니다. 변경으로 인해 AI가 적용대상이 되는 경우, 변경 주체의 의무가 자신의 변경으로부터 발생하는 위험으로 한정되도록 해 주시길 제언드리며, 변경 주체가 자신이 내리지 않았고 확인할 수도 없는 학습 단계의 결정에 대하여 책임을 부담하지 않도록 검토해 주시길 제언드립니다. 마찬가지로, 후속 변경으로 인해 발생하는 하위 단계의 위험은 변경 주체 자신의 사후학습 기법, 데이터셋 및 최종 용도의 선택에서 비롯되는 것으로서 개발자의 인지 및 통제 범위 밖에 있으므로, 기반 모델의 개발자에게 자동으로 귀속되지 않도록 검토해 주시길 제언드립니다.

제언: 고시(안) 제3조제1항제2호의 "실질적 변경"을 정의하고, 고시(안) 제3조제2항상 무엇이 "실질적 기여"에 해당하는지를 구체화해 주시길 제언드립니다. 아울러 제32조 의무가 AI 시스템이 아닌 AI 모델에 적용됨을 명확히 하고, 사업자가 AI 개발자인지 운전자(또는 이용사업자)인지에 따라 의무를 적절히 배분해 주시길 제언드립니다.

"위험" 및 "AI 관련 안전사고" 정의의 정교화 및 보고 기한의 연장

고시(안)은 인공지능사업자에 대하여 추가적인 사고 보고 의무를 신설하고 있으며, 이는 다음 조항에 규정되어 있습니다.

- **고시(안) 제2조제1호**은 "위험"을 “인공지능 수명주기 전반에 걸쳐, 사람의 생명·신체의 안전 또는 기본권이 침해될 가능성과 그 잠재적 피해의 심각성”으로 정의합니다.
- **고시(안) 제2조제4호**은 “인공지능 관련 안전사고”를 인공지능의 장애 등으로 인하여 위험이 현실적으로 발생하거나 공공의 안전에 중대한 위해가 발생한 경우로 정의합니다.
- **고시(안) 제8조제3항**은 해당 인공지능사업자에게 AI 관련 안전사고를 다음 3단계로 과학기술정보통신부에 보고하도록 요구합니다. (1) 사고발생 인지 시점으로부터 24시간 이내의 사고 발생 보고, (2) 사고발생 보고일로부터 7일 이내 초동조치 보고, (3) 사고발생 보고일로부터 15일 이내 사고처리 결과에 관한 보고입니다. 최종 보고서에는 사고의 원인, 피해 규모, 잔존 위험 제거를 포함한 후속 조치 및 재발 방지 계획이 포함되어야 합니다.

보고 의무를 촉발하는 AI 관련 안전사고가 생명, 신체의 안전, 기본권 또는 공공의 안전과 관련하여 높은 수준의 심각성 기준을 충족해야 한다는 점에는 분명합니다. BSA는 원칙적으로 높은 기준을 설정하는 데 동의합니다. 다만 다음과 같은 내용으로 제언 드립니다.

- **"위험"과 "AI 관련 안전사고"의 정의가 충분히 구체적이지 않습니다⁶.** 예를 들어, 사람의 "신체"에 대한 침해가 물리적 위해를 요구하는지(요구한다면 그 위해의 심각성 정도는 어떠한지)가 불분명합니다. 마찬가지로, 대한민국 헌법상 기본권으로 인정되는 범위가 광범위함을 고려할 때, 어떠한 사고가 보고 의무를 촉발할 만큼 심각한 침해에 해당하는지가 명확하지 않습니다. 보고 요건은 기업이 보고 대상 안전사고를 경미한 AI

⁶대한민국 헌법 제 10 조부터 제 37 조에 따르면, 기본권은 생명과 신체의 안전을 훨씬 넘어 평등(제 11 조), 직업선택의 자유(제 15 조), 사생활의 비밀과 자유(제 17 조), 표현의 자유 및 명예의 보호(제 21 조), 재산권(제 23 조) 등 광범위한 사항을 포괄합니다.

오류 및 오작동(예: 부정확한 출력)과 구별할 수 있도록, 객관적이고 합리적으로 확인 가능한 기준에 기초하도록 해 주시길 제언드립니다. 구체적인 기준이 없는 상황에서는 기업이 어떠한 사건이 보고 대상인지 판단하기 어렵습니다. 이는 과잉 보고로 이어질 수 있으며, 과잉 보고는 사고 해결에 투입되는 것이 더 바람직한 시간과 자원을 소모하는 한편, 공공의 안전을 실제로 위협하지 않는 다수의 신고 가운데서 실제로 위협이 되는 사고를 선별하도록 요구하게 됩니다. 나아가, 심각성이나 영향과 무관하게 현실화된 모든 위험에 대하여 사고 보고를 요구하는 것은 부정확한 출력과 같은 상대적으로 경미한 문제까지 사고로 보고되는 결과를 낳을 수 있습니다. 이는 인공지능사업자로 하여금 보고 의무를 회피하기 위하여 애초에 평가하는 위험의 범위를 축소하도록 유인할 수 있으며, 이는 과학기술정보통신부의 정책 목표에 역행하는 결과를 초래할 수 있습니다. 이러한 상황에서 사망, 중대한 신체적 상해 또는 상당한 금전적 손실 기준을 충족하는 재산 피해 등 높은 수준의 심각성 기준을 충족하는 구체적인 위험을 사고 보고의 기준으로 정의하는 것이 매우 중요할 것으로 판단합니다⁷.

- **보고 절차가 모호합니다.** AI 안전사고의 보고 절차는 상당한 실무적 어려움을 야기하는 측면이 있습니다. 특히 보고 요건이 AI 관련 안전사고의 발견을 전제로 하고 있고 그 개념 자체가 모호한 점이 우려됩니다. 기업은 사고 발생 여부를 신속하고 철저하게 조사해야 하는 경우가 많은데, 기업이 AI 관련 안전사고의 발생을 최초로 의심한 시점, 합리적으로 믿게 된 시점, 확인한 시점 중 어느 시점에 보고 요건이 촉발되는지가 불분명합니다. 이는 기초가 되는 정의의 불명확성으로 인하여 한층 가중되며, 그 결과 사업자는 보고 기한의 기산점뿐만 아니라 보고 의무 자체의 발생 여부에 대해서도 불확실한 상태에 놓이게 됩니다. 또한 고시(안)은 AI 가치사슬 전반에서 어떠한 주체가 보고 의무를 부담하는지를 명확히 하지 않고 있습니다. 개발자는 운전자(또는 이용사업자)가 이를 알려주지 않는 한 배포 과정에서 발생한 사고를 인지하지 못할 수 있으며, 어떠한 주체도 합리적으로 알 수 없었던 사고를 보고하지 않았다는 이유로 책임을 부과하는 것에 대해서는 다시 검토해 주시길 제언드립니다.
- **보고 기한이 짧습니다.** 최종 보고서 제출을 위한 15일의 기한 내에 기업은 사고의 원인을 규명하고, 피해를 정량화하며, 잔존 위험을 제거하고, 재발 방지 계획을 수립해야 합니다. 이러한 상세한 분석을 15일이라는 기간 내에 수행하는 것은 현실적으로 어려울 수 있습니다. 나아가 이러한 보고 요건은 사고가 진행 중인 상황에서, 기업이 시스템의 안전 확보에 투입해야 할 자원을 촉박한 보고 요건의 이행으로 전용하도록 이끌 수 있습니다. 그 대신, 최종 보고서는 고정된 기한을 설정하지 않고 실행 가능한 한 조속히 제출하도록 하는 것을 검토해 주시길 제언드립니다. 최초 보고에 대한 24시간의 기한 역시 상당히 더 긴 보고 기한을 두고 있는 유사 규제 체계(예: EU 인공지능법(EU AI Act), 캘리포니아 프런티어

⁷해외사례: EU 인공지능법 제 3 조 제 49 호는 "중대 사고(serious incident)"를 인공지능 시스템의 사고 또는 오작동으로서 직접 또는 간접적으로 다음 중 어느 하나를 초래하는 것으로 정의합니다: (a) 사람의 사망 또는 사람의 건강에 대한 중대한 위해, (b) 핵심 기반시설의 관리 또는 운영에 대한 중대하고 회복 불가능한 중단, (c) 기본권 보호를 목적으로 하는 EU 법상 의무의 침해, (d) 재산 또는 환경에 대한 중대한 위해.

캘리포니아 프런티어 인공지능 투명성법도 유사한 접근방식을 취하여, "치명적 안전사고(critical safety incident)"를 다음 중 어느 하나로 정의합니다: (a) 파운데이션 모델(foundation model)의 모델 가중치에 대한 무단 접근, 변경 또는 유출로서 사망, 신체 상해 또는 재산의 손상·손실을 초래하는 것, (b) 재앙적 위험(catastrophic risk)의 현실화로 인하여 발생하는 위해, (c) 사망 또는 신체 상해를 야기하는 파운데이션 모델에 대한 통제력 상실, (d) 파운데이션 모델이 해당 행동을 유도하기 위해 설계된 평가의 맥락 외에서, 개발자의 통제 또는 모니터링을 회피하기 위하여 개발자를 상대로 기만적 기법을 사용하고, 그 방식이 재앙적 위험의 실질적 증가를 보여주는 경우.

인공지능 투명성법(Transparency in Frontier AI Act))와 정합성이 부족한 측면이 있습니다.⁸.

- **기존 보고 체계와의 정합성이 확보되어 있지 않습니다.** 고시(안)이 사고 발생 이후 인공지능사업자에게 잔존 위험의 해소 또는 재발 방지를 요구하는 경우, 사업자가 NIST 사이버보안 프레임워크(NIST Cybersecurity Framework)나 ISO/IEC 27035 등에 기반한 기존의 사고 대응 체계에 부합하는 방식으로 해당 요건을 충족할 수 있도록 명시적으로 허용하는 방안을 검토할 필요가 있습니다. 다수의 AI 안전사고, 특히 보안 취약점이나 악의적 악용과 관련된 사고는 이미 사업자의 사이버보안 사고 대응 프로그램의 범위에 포함되어 있으며, 중복적인 AI 특화 절차는 안전 성과의 개선 없이 컴플라이언스 부담만 가중시킬 것입니다.

제언: 다음 사항을 검토해 주실 것을 제언드립니다.

- "위험"과 "AI 관련 안전사고"를 사망, 중대한 신체적 상해 또는 중대한 재산 피해와 같이 높은 수준의 심각성 기준을 충족하는 구체적인 위해를 기준으로 정의해주시길 제언드립니다.
- 인공지능사업자가 보고 대상 사고의 발생을 확인한 경우에 한하여 보고 의무가 촉발됨을 명확히 해주시길 제언드립니다.
- EU 인공지능법, 캘리포니아 프런티어 인공지능 투명성법 등 유사 규제 체계의 보고 기한과 정합성을 확보할 수 있도록 최초 보고 기한을 연장하는 것을 고려해주시길 제언드립니다.
- 어떠한 사업자도 합리적으로 알 수 없었던 사고에 대한 보고 책임을 지지 않도록, 보고 의무를 부담하는 주체를 명확히 하는 것을 제언드립니다.
- 의미 있는 근본 원인 분석을 완료하기에 충분하지 않을 수 있는 고정된 기한이 아니라, 실행 가능한 한 조속히 최종 보고서를 제출할 수 있도록 정책 수립을 제언드립니다.
- 사업자가 기존의 사고 대응 체계에 부합하는 방식으로 잔존 위험 및 재발 방지 관련 요건을 충족할 수 있음을 명확히 해주시길 제언드립니다.

마치며

고시(안)에 대하여 의견을 제시할 기회를 주셔서 진심으로 감사드립니다. BSA의 제언이 인공지능기본법의 개선 논의에 도움이 되기를 바랍니다. BSA는 인공지능기본법 제도 개선 연구반에 참여하게 되어 진심으로 기쁘게 생각하며, 협력을 계속해 나가기를 기대합니다. 본 의견서와 관련하여 문의 사항이 있으시거나 추가로 필요하신 부분이 있으면 말씀 부탁드립니다.

감사합니다.

Tham Shen Hong
아시아·태평양 정책 이사

⁸EU 인공지능법 제 73 조에 따르면, 고위험 인공지능 시스템과 관련된 중대 사고는 15 일 이내에 보고되어야 하며, 사망을 야기하였을 수 있는 경우에는 10 일로, 광범위한 침해 또는 핵심 기반시설의 중대하고 회복 불가능한 중단의 경우에는 2 일로 단축됩니다. 캘리포니아 프런티어 인공지능 투명성법에 따르면, 치명적 안전사고는 발견 후 15 일 이내에 보고되어야 하며, 사망 또는 중대한 신체적 상해의 임박한 위험을 야기하는 경우에 한하여 24 시간 이내에 보고되어야 합니다.