

AI시대의 일하는 방식 변화와 HRD 전략

김동호 (서울대학교 산업인력개발학과 교수)

I. 서론

대규모 언어모델(LLM)을 포함한 생성형 AI는 더 이상 미래 기술이 아니라, 이미 일부 지식노동 현장에서 측정 가능한 성과를 내고 있다. 예컨대 중간 수준의 전문직 글쓰기 과업을 대상으로 한 실험연구에서는 ChatGPT 사용 집단의 과업 수행 시간이 평균 40% 감소하고 산출물 품질이 18% 향상되었으며, 실제 고객지원 업무 현장을 분석한 연구에서는 생성형 AI 지원 도구가 시간당 해결 건수를 평균 15% 높이는 효과를 보였다(Noy & Zhang, 2023). 특히 후자의 연구는 이러한 효과가 모든 근로자에게 동일하게 나타나는 것이 아니라, 초보자와 저숙련 근로자에게 더 크게 나타난다는 점도 함께 보여준다. 이는 생성형 AI의 효과를 단순한 “생산성 향상”으로 일반화하기보다, 과업의 구조와 숙련 분포에 따라 상이하게 작동하는 현상으로 보아야 함을 시사한다(Brynjolfsson et al., 2025).

그러나 이러한 변화가 곧바로 직업 전체의 일괄적 전환이나 소멸을 의미하는 것은 아니다. 기술 변화의 노동시장 효과를 설명하는 과업 혹은 스킬 기반 연구들은 기술이 대체하는 단위가 직업(job) 전체가 아니라 직업을 구성하는 개별 과업(task)이나 스킬(skill) 단위임을 강조한다. 자동화는 명시적 규칙으로 수행 가능한 과업을 대체하는 반면, 비반복적 문제 해결과 복합적 의사소통 과업은 오히려 인간의 수행을 보완하는 경향이 있다(Autor et al., 2003). 동일한 직무 내에서도 과업 구성은 상당히 이질적이기 때문에 AI활용 양상에 따른 성과와 수행은 크게 달라질 수 있다. 더 나아가 자동화는 일부 과업을 대체하지만, 동시에 인간에게 비교우위가 있는 새로운 과업을 재창출하기도 한다. 따라서 AI 시대 노동 변화를 이해하는 핵심 분석 단위는 “직업명”이 아니라 과업 묶음의 재구성 방식으로 보아야 한다(Autor &

Handel, 2013; Acemoglu & Restrepo, 2019).

이 때문에 조직의 핵심 과제는 “AI를 도입할 것인가”가 아니라 “AI와 인간의 기여를 어떻게 다시 설계할 것인가”로 이동한다(Parker & Grote, 2022). 인간-AI 협력, 하이브리드 지능(hybrid intelligence), 그리고 디지털 시대의 직무재설계 등의 연구는 조직 성과가 기술 그 자체보다 인간의 판단과 AI의 계산 역량을 어떤 구조로 연결하느냐에 더 크게 좌우된다고 본다(Jarrahi, 2018). 동시에 AI 도입은 데이터·거버넌스·전략을 포함하는 조직적 준비도를 필요로 하며, HRD는 더 이상 구성원들의 테크놀로지 사용법을 교육시키는 보조 기능에 머물 수 없다(Jöhnk et al., 2021; Li & Yeo, 2024). 오히려 HRD는 근로자가 AI와 협업하면서도 주도성과 전문성을 잃지 않도록 사회기술적 조건을 설계하는 핵심 기능으로 재정의되어야 한다(Dellermann et al., 2019).

II. 일터에서 직무 수행의 양상 변화와 직무 재설계

1. 직무의 해체와 과업 기반 접근

AI가 노동에 미치는 영향을 정확히 진단하기 위해서는 거시적 직업 분류보다 과업기반 접근이 더 유효할 것이다. 하나의 직업은 분석, 자료 처리, 커뮤니케이션, 판단, 관계 조정, 예외 처리와 같은 이질적 과업들의 결합으로 이루어져 있으며, 같은 직무 안에서도 실제 수행 과업은 개인과 조직에 따라 상당히 달라질 수 있다(Autor & Handel, 2013). 따라서 직업 단위의 평균적 전망만으로는 AI 노출도나 직무 재설계의 필요성을 충분히 설명할 수 없고, HRD 역시 역량 사전을 직업명 중심으로만 관리해서는 변화의 실재를 포착하기 어렵다(Autor et al., 2003).

고전적인 과업기반 연구가 보여주듯, 디지털기술은 규칙 기반의 반복 과업을 주로 대체해 왔다. 그러나 생성형 AI는 이러한 대체 범위를 언어 기반의 일부 인지 과업으로 확장시키고 있다(Noy & Zhang, 2023). 현재까지의 실증 근거가 가장 분명한 영역도 바로 이런 종류의 과업들이다. 즉, 초안 작성, 요약, 패턴화된 고객 응대, 규칙이 비교적 명확한 텍스트 생산처럼 언어적 형식은 복잡하지만 맥락은 비교적 제한된 과업에서는 AI의 성과가 비교적 잘 관찰된다. 반면 복합적 맥락 판단, 관계 조정, 윤리적 책임성, 불완전한 정보 속의 문제정의와 같은 과업에서는 인간의 역할이 여전히 핵심적이다(Jarrahi, 2018). 따라서 AI는 지식노동 전체를 일괄적으로

대체한다기보다, 특정 과업을 자동화하고 다른 과업에서는 인간 역량을 증강하는 방향으로 작동한다(Brynjolfsson et al., 2025; Dellermann et al., 2019).

이러한 점에서 HRD 담당자들과 연구자들이 해야 할 일은 단순히 “AI를 쓸 수 있는가”를 묻는 것이 아니라, 어떤 과업은 자동화하고, 어떤 과업은 증강하며, 어떤 과업은 인간이 계속 주도해야 하는가를 분리해 판단하는 것이다(Parker & Grote, 2022). 이 구분은 기술 분류가 아니라 책임, 학습 기회, 성과지표, 그리고 장기적 전문성 형성의 문제로 바라봐야 한다. 디지털 기술 환경에서 직무재설계에 주목하는 연구들은 바로 이 지점을 강조한다. 설계가 잘되면 AI는 자율성·피드백·학습 가능성을 높일 수 있지만, 설계가 나쁘면 감시·세분화·탈숙련화로 이어질 수도 있다(Tambe et al., 2019).

싱가포르의 리관유 혁신 도시 센터가 발표한 AI 시대의 직무 재설계 가이드는 이와 같은 과업 기반 분석을 실제 조직의 인적자원관리에 적용한 모범적인 사례로 꼽힌다(GPAI, 2025). 이 가이드는 단순히 비용 절감을 위한 기술 도입을 경계하며, 기업의 비즈니스적 요구를 충족시키는 동시에 근로자의 기여 가치를 향상시키고 직장 내의 신뢰를 유지하는 인간 중심적 접근을 강조한다(<표 1> 참조).

<표 1> 조직 AI 내재화를 위한 직무 재설계 프로세스

프로세스 단계	명칭	세부 수행 내용 및 인적자원개발 관점의 의미
1단계	직무 분해	현재 직원이 수행하고 있는 직무를 분석하여 세부적인 개별 과업 단위로 분해한다.
2단계	잠재적 영향 평가	분해된 각각의 과업에 대하여 AI 기술이 미칠 수 있는 대체 가능성 및 증강 효과를 정량적, 정성적으로 평가한다.
3단계	기술 도입 수준 결정	AI 기술의 현재 성숙도를 고려하여 각 과업에 알고리즘을 어느 정도의 범위와 깊이로 적용할 것인지 실무적 타당성을 검토한다.
4단계	가치 평가 및 직원 협의	조직의 관리자와 실제 과업을 수행하는 직원이 참여하여, 인간 고유의 통찰이나 대인 관계 능력이 요구되는 핵심 과업이 무엇인지 논의하고 보존 가치를 결정한다.
5단계	이행 계획 수립	자동화 및 증강으로 인해 절감된 시간을 고부가가치 과업으로 재할당하기 위한 타임라인과 필요 학습 목표를 구체적으로 수립한다.
6단계	미래형 직무 재조립	분석과 타협의 과정을 거쳐 AI와 인간이 유기적으로 협력할 수 있는 미래 지향적이고 고차원적인 새로운 역할로 과업들을 재조립한다.

인간 중심 직무 재설계는 비용 절감이나 인력 대체의 목적으로 출발해서는 안 된다. 보다 적절한 출발점은 생산성, 판단의 질, 근로자 발달을 동시에 최적화하는 사회기술적 설계다. 이러한 관점에서 Parker와 Grote (2022)는 디지털 기술이 일의 구조를 바꾸는 방식이 단순히 효율성만의 문제가 아니라고 지적한다. 같은 기술도 자율성과 의미를 확대할 수도 있고, 반대로 작업을 단절시키고 통제 강도를 높일 수도 있기 때문이다. 경영학에서도 AI는 자동화와 증강이라는 상반되지만 상호의존적인 두 논리를 동시에 내포한 “automation-augmentation paradox”로 설명된다. 핵심은 기술의 성능이 아니라, 권한과 책임, 개입 시점, 인간 판단의 위치를 어떻게 배치하는가이다(Raisch & Krakowski, 2021).

또한 직무 재설계의 지속 가능성은 구성원 참여에 크게 좌우된다. 선행연구들은 변화 목표와 개입 방식의 설계에 조직 구성원들이 참여할 때 수용성과 실행 가능성이 높아지며, 참여적 설계를 통한 근로자의 심리적 준비도 향상과 실제 참여가 핵심 변수임을 보여준다(Ullrich et al., 2023). 다시 말해, 조직 구성원들은 새로운 도구의 수동적 사용자라기보다 미래 직무의 공동 설계자여야 한다. HRD는 교육 이후의 적응을 기다리기보다, 도입 초기부터 직무 해체·과업 재배치·검증 규칙 설계 과정에 직원이 참여하도록 만드는 구조를 마련해야 하는 것이다(Orso et al., 2022; Abhari, 2025).

2. 지식 노동의 새로운 표준: 인간-AI협업 역량과 조직 대응 전략

생성형 AI가 확산될수록 지식노동의 가치 분포도 달라진다. 텍스트 생성이나 요약과 같은 표준화된 응답 생성의 비용은 급격히 낮아지기 때문에, 주목받아야 할 역량은 백지에서 무엇을 처음부터 만들어내는 능력이 아니다. 오히려 시급한 문제를 포착하는 능력, AI에게 위임할 과업을 결정하는 능력, 그리고 AI와 협업하여 도출한 결과물을 보완하고 현업에 맥락화하는 능력이다. 이는 생성형 AI가 초안 생성과 같은 경계가 분명한 과업에서 시간 비용을 크게 줄이고 품질을 높인다는 실증 결과와도 부합한다. 다시 말해, 지식노동의 중심이 “생산”에만 있는 것이 아니라, “생산 이후의 최적화”로 무게중심을 분배한다고 봐야 한다(Noy & Zhang, 2023; Brynjolfsson et al., 2025).

최소한 지금 시점에서는, AI로 인한 자동화가 단순 작업에 대한 인간의 완전 해방을 의미하는 것은 아니다. AI가 최초로 산출해낸 작업물에 대한 철저한 검토와

정교화가 필요하고 이는 결국 인간의 영역 지식과 전문성의 발휘를 요구한다. 자연어 생성 연구에서 잘 알려져 있듯, 거대생성모델은 그럴듯하지만 근거 없는 내용을 산출하는 환각(hallucination)에 여전히 취약하다. 더 나아가 사용자는 맥락적 단서가 의심을 요구하는 상황에서도 AI의 조언을 과도하게 따를 수 있으며, 이 과잉의존은 조직의 손실과 실패를 초래할 수도 있다. 또한, 거대 생성 모델이 스스로 학습한 데이터에 과하게 의존한다거나 접근이 허용된 자료에 우선적인 가중치를 주기 쉽다. 결국 문제는 단순한 정확도 부족이 아니라, AI가 제시한 결과를 인간이 얼마나 정밀하게 보정하는가의 문제다. 즉, 언제 AI를 신뢰하고, 언제 검증하고, 언제 인간이 개입해 보완해야 하는가가 핵심이 된다(Ji et al., 2023; Klingbeil et al., 2024; Glikson & Woolley, 2020)

정리하자면, AI 시대의 지식노동은 면밀한 검증·보완 중심의 워크플로우로 재정의되어야 한다. 이때 검증이 의미하는 바는 단순히 AI의 최초 산출물의 신뢰도를 판단하는 것을 넘어선다. AI에게 위임할 영역에 대한 의사결정, 산출물의 형태에 대한 구상, 이상적 산출물을 도출할 때까지의 순환적 협업과 프롬프팅, 출처 대조, 예외 처리, 감사 가능성과 같은 포괄적 과업들을 포함한다. 인간-AI 협업이 성과를 내기 위해서는 위임이 직관이나 즉흥성에 맡겨져서는 안 되며, 어떤 상황에서 인간이 재개입해야 하는지에 대한 메타지식이 필요하다. 그러므로 HRD의 과제는 AI 사용법 교육에 머무르지 않고, 포괄적이고 깊이있는 검증 절차를 조직의 훈련 가능한 역량으로 제도화하는 데 있다(Fugener et al., 2022; Li & Yeo, 2024; Tambe et al., 2019).

AI 기반 직무 재설계의 핵심은 인간과 AI의 대체 관계가 아니라 보완 관계를 전제로 한 과업 재배치에 있다. 조직 맥락에서 AI는 대규모 데이터 처리, 패턴 인식, 예측, 반복적 연산과 같은 과업에서 강점을 보이는 반면, 인간은 문제의 정의, 예외 상황의 판단, 맥락 해석, 윤리적 숙고, 대인 조정과 같은 비정형 과업에서 상대적 우위를 유지한다(Jarrahi, 2018; Dellermann et al., 2019; Parker & Grote, 2022). 따라서 AI 도입의 실질적 성과는 “얼마나 많은 업무를 자동화했는가” 보다 “어떤 과업을 누구에게 재배분했는가”에 의해 좌우된다. 이런 점에서 AI는 단순한 도구를 넘어, 협업 과정에 참여하는 가상의 팀원 혹은 팀 성과를 좌우하는 사회기술적 행위자로 이해될 수 있다(Seeber et al., 2020).

특히, 인간-AI 협업의 성패를 가르는 가장 중요한 매개는 여전히 신뢰의 형성이다. 기존 연구는 인간이 AI를 인간 동료와 동일한 방식으로 신뢰하지 않으며,

특히 초기 협업 상황에서는 정서적·관계적 신뢰보다 성능, 정확성, 일관성에 기반한 인지적 신뢰를 먼저 형성하는 경향이 있음을 보여준다(Glikson & Woolley, 2020; Gillath et al., 2021). 또한 인간-AI 팀에서는 팀의 규모, 상호작용 구조, AI의 역할 정의에 따라 신뢰의 양상이 달라진다. 예컨대 소규모 협업에서는 인간-AI 조합이 인간-인간 조합보다 낮은 상호 신뢰로 이어지지만, 과업 구조가 명확하고 역할이 적절히 설계된 경우 이러한 차이는 완화될 수 있다(Georganta & Ulfert, 2024). 따라서 협업에서의 과제는 AI에 대한 무조건적 신뢰를 조성하는 것이 아니라, 설명 가능성, 검토 가능성, 인간의 최종 판단권을 보장함으로써 ‘교정된 신뢰(calibrated trust)’를 형성하도록 돕는 데 있다.

나아가 인간-AI 협업은 단순한 효율화가 아니라 인지적 여유의 재배분이라는 관점에서 이해될 필요가 있다. 반복 업무와 정보 탐색 부담이 감소할수록 근로자는 문제 재정의, 대안 비교, 창의적 조정, 관계 관리와 같은 상위 수준의 과업에 더 많은 주의를 배분할 수 있다(Parker & Grote, 2022). 조직 수준에서 AI 수용역량은 조직 창의성과 성과에 긍정적 연관성을 보이며(Mikalef & Gupta, 2021), 심리적 안전감이 확보된 팀일수록 실험, 질문, 오류 공유를 통한 학습 행동이 촉진된다(Edmondson, 1999). 따라서 우리는 AI를 감시와 통제의 장치가 아니라 역량 증강의 파트너로 경험하게 만드는 직무 설계, 피드백 체계, 팀 규범을 함께 구축해야 한다.

Ⅲ. AI 내재화를 위한 조직 재구성 패러다임

AI의 심층 도입은 개별 직무 수준의 자동화를 넘어, 데이터 인프라, 의사결정 구조, 역할 분담, 학습 체계를 동시에 재구성하는 사회기술적 변화다. 관련 연구들도 조직의 AI 도입 성패가 모델 성능 자체보다 AI 준비도, 데이터 거버넌스, 프로세스 통합, 조직학습 역량에 더 크게 좌우된다고 본다(Holmstrom, 2022; Johnk et al., 2021; Mikalef & Gupta, 2021; Trocin et al., 2021). 따라서 AI 내재화는 도구의 추가가 아니라 운영모델의 재설계로 이해되어야 한다.

1. 조직 내 AI 도입의 3단계 성숙도 모델

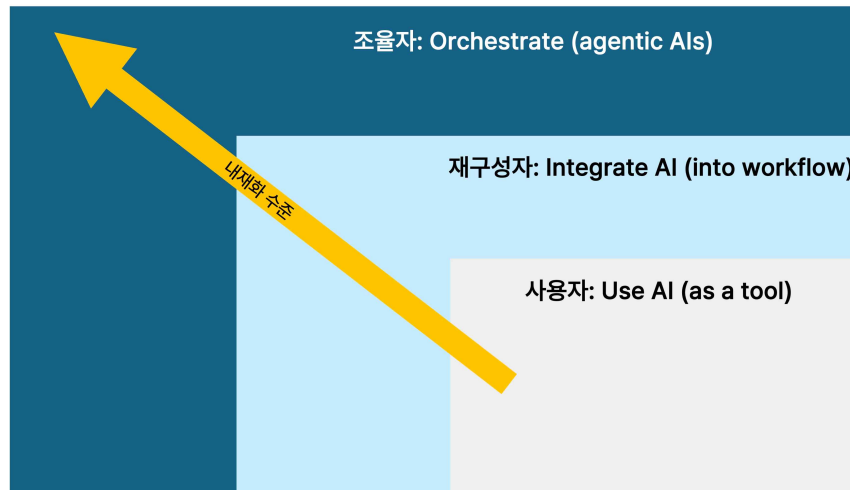
본 이슈페이퍼에서는 조직내 AI 활용과 관련된 선행연구들 (Holmstrom, 2022; Johnk et al., 2021; Mikalef & Gupta, 2021; Trocin et al., 2021; Yalenios &

d’ Armagnac, 2023)을 종합하여 AI 내재화와 관련된 조직의 성숙도와 HRD가 주목해야 할 관련 요소를 <표 2>에 보듯이 세 단계로 제시하고자 한다.

<표 2> 조직의 AI 내재화를 위한 성숙도 수준

단계	성숙도 수준	주요 특징	HRD의 초점
1단계 사용자	기반 구축 및 개별적 사용 (Use)	개인 또는 일부 부서가 범용 AI 도구를 기본 실험적으로 사용. 활용은 산발적이며, 데이터 · 보안 · 윤리 기준이 불완전함. 형성	AI 리터러시, 사용 가이드라인, 초기 수용성
2단계 재구성자 (Integrate)	워크플로우 통합	AI가 업무 프로세스와 데이터 흐름에 연결되어 표준 업무 일부를 재구성. 성과 측정, 역할 재설계, 부서 간 협업이 중요해짐.	직무 재설계, 프로세스 학습, 검증 루틴, 부서 간 협업 역량
3단계 조율자	하이브리드 지능 오케스트레이션 (Orchestrate)	복수의 AI 에이전트를 인간 전문가가 하나의 운영 체계 안에서 동적으로 조율함. 성과는 개별 도구보다 조율 품질에 좌우됨.	시스템 사고, 예외 처리, 거버넌스, 윤리 · 위험 관리, 조율 역량

이 구분의 핵심은 조직이 AI를 “도구로 사용” 하는 수준에서 “업무 흐름에 통합” 하고, 나아가 “하이브리드 지능을 조율” 하는 수준으로 이동할수록, 기술보다 조직 설계와 학습 체계의 중요성이 커진다는 점이다. 즉, 성숙한 AI 조직의 경쟁우위는 개별 모델의 성능보다 사람 · 데이터 · 프로세스 · 거버넌스를 얼마나 정합적으로 연결하느냐에 의해 결정된다. 주의해야 할 점은, 이 3단계의 내재화 단계가 상호배타적이라기 보다는 상위의 단계가 하위 단계를 포함하는 확장적 성격이 강하다는 것이다. 예를 들어, 다양한 AI를 오케스트레이트 할 수 있는 구성원은 개별적 AI를 사용 및 워크플로우에 통합을 하는 성격을 보인다. 반면, AI를 워크플로우에 통합을 한다고 해서 반드시 AI 에이전트 오케스트레이션 단계에 이르지 않는다. 이러한 매커니즘은 아래의 [그림 1]처럼 표현될 수 있다.



[그림 1] AI 내재화 수준과 특성 관계도

2. 다중 AI 환경에서의 오케스트레이터 역량 개발

오케스트레이션 단계에 도달한 조직에서는 관리자와 지식근로자의 핵심 역할이 더 이상 직접 수행 자체에 머무르지 않는다. 중요한 것은 누가 더 많은 업무를 손으로 처리하느냐가 아니라, 여러 AI 기능과 인간의 전문성을 어떻게 결합하여 더 나은 결과를 산출하느냐이다. 이 과정에서 직무의 중심은 실행에서 조율로, 생산에서 검증으로, 개인 수행 능력에서 사회기술적 설계 능력으로 이동한다. 특히 하이브리드 지능(hybrid intelligence) 관점에서 볼 때 조직의 성과는 인간과 AI 가운데 어느 한쪽의 우수성에 의해 결정되기보다, 양자의 상호보완적 결합 구조가 얼마나 정교하게 설계되어 있는가에 의해 좌우된다(Dellermann et al., 2019). 또한 인간-AI 팀 협업 연구는 AI가 단순한 도구가 아니라 정보 제안, 대안 생성, 의사결정 지원에 참여하는 준행위자적 성격을 띠 때, 전통적인 상명하달식 관리보다 새로운 형태의 협업 규칙과 조율 메커니즘이 요구된다고 지적한다(Seeber et al., 2020).

이러한 맥락에서 오케스트레이터에게 요구되는 첫 번째 역량은 문제 프레임링 역량이다. 많은 조직이 AI 도입에 실패하는 이유는 기술 성능이 낮아서라기보다, 애초에 문제를 제대로 정의하지 못하기 때문이다. 즉, 어떤 업무가 데이터 처리와 속도에 의해 좌우되는지, 어떤 업무가 맥락 판단과 책임 있는 해석을 요구하는지 구분하지 못한 채 AI를 만능 해결책처럼 투입할 경우, 성과는 불안정해지고 현장 수용성도 낮아진다. 따라서 오케스트레이터는 업무를 표면적 절차가 아니라 판단 단위와 불확실성 수준에 따라 재구성해야 하며, 어느 지점까지는 AI의 제안을

활용하고 어느 지점부터는 인간의 숙고와 승인 절차가 반드시 개입해야 하는지를 설계해야 한다. 이는 단순한 기술 이해가 아니라 업무 구조, 책임 흐름, 위험 지점을 동시에 파악하는 문제 구조화 능력을 뜻한다(Raisch & Krakowski, 2021; Parker & Grote, 2022).

두 번째 핵심 역량은 판단·검증 역량이다. 다중 AI 환경에서는 서로 다른 모델이나 도구가 각기 다른 데이터셋, 알고리즘적 전제, 추론 방식을 바탕으로 결과를 생성한다. 따라서 오케스트레이터는 단일 산출물을 소비하는 사용자가 아니라, 복수의 결과를 비교·대조하고 품질을 점검하는 검토자로 기능해야 한다. 여기에는 데이터의 적합성, 입력값의 편향 가능성, 모델 출력의 논리적 정합성, 설명 가능성, 맥락적 타당성을 확인하는 능력이 포함된다. 특히 생성형 AI는 그럴듯하지만 부정확한 응답을 생성할 수 있기 때문에, 결과가 얼마나 유창하게 보이느냐보다 그것이 실제 문제 상황에 적합한가를 따지는 검증 절차가 필수적이다. 인간-AI 위임 연구 역시 효과적인 협업은 “위임 자체”보다 “언제 위임을 멈추고 인간이 재개입하는가”에 달려 있다고 보고한다(Fugener et al., 2022). 결국 오케스트레이터는 기술의 속도를 추종하는 사람이 아니라, 속도와 정확성, 효율과 책임 사이의 균형점을 찾아내는 판단자여야 한다.

세 번째로 중요한 것은 프로세스 설계 역량이다. 다중 AI 환경에서는 개별 도구를 잘 쓰는 능력만으로는 충분하지 않다. 더 중요한 것은 여러 도구와 인간 전문가의 기여를 하나의 흐름으로 연결하는 능력이다. 예를 들어 어떤 과업에서는 생성형 AI가 초안을 만들고, 다른 분석 도구가 수치 검증을 수행하며, 인간 전문가가 예외 상황을 판단하고, 마지막으로 관리자가 윤리적·전략적 승인 여부를 결정할 수 있다. 이처럼 실제 업무는 선형적이기보다 다층적이며, 오케스트레이터는 이 과정에서 작업 순서, 검토 지점, 피드백 환류, 오류 발생 시의 재처리 절차를 설계해야 한다. 이는 곧 업무 흐름을 “누가 일하는가”의 차원에서 보는 것이 아니라, “어떤 정보가 어떤 시점에 어떤 기준으로 이동하고 점검되는가”의 차원에서 보는 시스템적 사고를 요구한다(Johnk et al., 2021). 따라서 HRD는 도구별 기능 교육을 넘어, 프로세스 맵핑, 예외 처리 루틴 설계, 결과 검토 기록화와 같은 고차적 설계 역량을 훈련해야 한다.

네 번째는 거버넌스 역량이다. AI 활용이 조직 전반으로 확산될수록 효율성 못지않게 중요한 것이 책임 소재, 윤리성, 공정성, 그리고 인간의 최종 통제권이다. 특히 다중 AI 환경에서는 오류가 어느 지점에서 발생했는지, 편향이 입력 데이터에서 비롯된 것인지 모델 구조에서 비롯된 것인지, 혹은 인간의 검토 부재 때문인지가

불분명해질 수 있다. 이때 오케스트레이터는 단순히 성과를 추적하는 관리자라기보다, 의사결정 구조를 설계하고 책임의 경계를 명료화하는 메타 관리자여야 한다. 인간-AI 협업 연구는 신뢰가 높은 조직일수록 AI 활용도가 높아지는 것이 아니라, 오히려 설명 가능성과 인간의 감독 가능성이 확보될 때 비로소 지속 가능한 신뢰가 형성된다고 본다(Seeber et al., 2020; Parker & Grote, 2022). 따라서 오케스트레이션 역량에는 성과 향상을 위한 조율뿐 아니라, 시스템이 어떤 규칙 아래 작동하는지, 누가 최종 판단을 내리는지, 이의 제기와 수정은 어떻게 가능한지를 구조화하는 책임 있는 설계가 포함되어야 한다.

3. 변화 대응력과 민첩성의 중요성

전통적 조직 설계는 오랫동안 비용 절감, 표준화, 예측 가능성, 통제 가능성을 중심 원리로 삼아 왔다. 이러한 원리는 대량생산과 안정적 수요를 전제로 한 산업화 시대에는 매우 효과적이었다. 그러나 최근 조직이 직면하는 환경은 과거와 본질적으로 다르다. AI가 촉발한 산업 현장의 실시간 변화는 고정된 프로세스와 사전 계획만으로는 대응하기 어려운 불확실성을 만들어내고 있다. 이처럼 환경 변화의 빈도와 강도가 높아질수록, 기존의 효율성 중심 설계는 오히려 취약성을 심화시킬 수 있다. 효율성 극대화를 위해 여유 자원과 중복 구조를 제거한 조직일수록 충격 발생 시 적응력이 약해지기 때문이다. 이러한 배경 속에서 조직의 경쟁우위는 동일한 자원을 더 싸게 운영하는 능력보다, 변화의 신호를 빨리 감지하고 이를 해석하며 구조를 빠르게 재조정하는 변화 대응력으로 이동하고 있다(Holmstrom, 2022; Overby et al., 2006).

이러한 논의는 전략경영의 역동적 역량(dynamic capabilities) 관점과도 맞닿아 있다. 해당 이론은 조직의 지속적 경쟁우위가 기존 자원을 효율적으로 활용하는 능력만으로는 설명되지 않으며, 환경 변화 속에서 새로운 기회를 감지하고(sensing), 이를 포착하며(seizing), 기존 자원과 구조를 재구성(reconfiguring)하는 능력에서 비롯된다고 본다(Teece et al., 1997). AI는 바로 이 세 과정에 모두 개입할 수 있는 기술이다. 예컨대 실시간 데이터 처리와 이상 탐지를 통해 변화의 조짐을 더 빨리 감지할 수 있고, 다양한 시나리오를 시뮬레이션함으로써 의사결정 대안을 더 정교하게 비교할 수 있으며, 운영 과정에서 축적되는 데이터를 바탕으로 프로세스를 지속적으로 재조정할 수 있다. 따라서 AI의 조직적 가치는 단순히 사람의 일을

대신하는 자동화에 있지 않고, 조직이 불확실성을 해석하고 대응하는 능력을 증폭시키는 데 있다.

이러한 전환을 잘 보여주는 대표적 사례가 디지털 공급망 트윈(digital supply chain twin)이다. 디지털 트윈은 물리적 시스템의 상태를 데이터 공간 안에 거의 실시간으로 반영하고, 이를 기반으로 예측과 시뮬레이션을 수행하는 디지털 복제 구조를 의미한다. 공급망 맥락에서 이는 재고, 운송, 수요, 생산 상태 등 복합적 요소를 통합적으로 가시화하고, 특정 교란 상황이 발생했을 때 그 파급 경로와 대안의 효과를 사전에 실험해볼 수 있게 한다(Ivanov & Dolgui, 2021). 중요한 것은 여기서 AI가 단지 예측 정확도를 높이는 분석 도구에 머무르지 않는다는 점이다. 오히려 AI는 조직이 다양한 불확실성 하에서 “무엇이 일어날 수 있는가”를 가정하고 “그때 어떤 선택이 가장 바람직한가”를 비교하게 만드는 시뮬레이션 기반 사고의 토대를 제공한다. 이는 문제 발생 이후의 사후 대응을 넘어, 잠재적 위험을 선제적으로 관리하는 조직 역량을 가능하게 한다.

하지만 민첩성은 기술이 도입되었다고 자동으로 생겨나는 속성이 아니다. 디지털 전환 연구는 조직이 데이터를 더 많이 보유하고 고도화된 분석도구를 도입하더라도, 그 결과를 해석하고 행동으로 연결할 수 있는 구조가 없으면 성과가 제한적이라고 지적한다(Holmstrom, 2022). 다시 말해 민첩성은 기술적 자동화의 부산물이 아니라, 기술·조직·인간 역량이 정렬될 때 나타나는 결과이다. 예측 알고리즘이 위험 신호를 감지하더라도 현장 리더가 그것을 신뢰하지 않거나, 부서 간 정보 공유가 차단되어 있거나, 의사결정 권한이 지나치게 상층부에 집중되어 있으면 조직은 여전히 느리게 반응할 수밖에 없다. 따라서 민첩성은 분석 도구의 성능보다, 그 결과가 얼마나 신속하게 공유되고 해석되며 실행되는가에 의해 결정된다. 이런 점에서 AI 시대 조직의 경쟁력은 ‘더 똑똑한 모델’을 갖는 것보다 ‘더 빠르게 학습하고 조정하는 조직’을 만드는 데 있다.

HRD의 관점에서 보면, 여기서 가장 중요한 과제는 민첩성을 단순히 운영 부서나 디지털 부서의 전문 역량으로 한정하지 않는 것이다. AI 기반 민첩성이 실질적인 조직 성과로 연결되기 위해서는 데이터 해석 능력, 시나리오 기반 사고, 교차기능 협업, 예외 상황에서의 학습, 빠른 피드백 수용과 같은 역량이 조직 전반에 내재화되어야 한다. 이는 개인 차원의 기술 습득을 넘어 집단 차원의 학습 체계와 연결된다. 예컨대 구성원들이 변화 신호를 감지했을 때 이를 보고하고 공유하는 문화, 가설적 시나리오를 빠르게 검토하는 회의 방식, 실패 사례를 비난보다 학습 자원으로

전환하는 성찰 구조가 함께 작동해야 한다. 결국 민첩성은 단순히 빠르게 움직이는 것이 아니라, 빠르게 감지하고, 빠르게 해석하고, 빠르게 재구성하는 학습 체계를 의미한다(Overby et al., 2006).

IV. AI시대 HRD의 설계원리와 실천방향

이상의 논의를 종합하면, AI 시대 HRD의 과제는 더 이상 교육 콘텐츠를 공급하고 정기 교육을 운영하는 기능적 부서에 머무르지 않는다. 오히려 HRD는 인간과 기술이 함께 성장할 수 있는 사회기술적 학습 환경의 설계자로 전환해야 한다. 생성형 AI와 다양한 자동화 도구가 업무 전반으로 확산될수록, 조직 내 학습은 더 이상 특정 시점에 별도로 분리된 활동으로 유지되기 어렵다. 학습은 일과 병렬적으로 존재하는 부가 활동이 아니라, 업무 수행 과정 속에서 끊임없이 발생하고 갱신되는 내재적 과정이 되어야 한다. 이때 HRD의 핵심 역할은 단순히 새로운 기술의 사용법을 안내하는 것이 아니라, 구성원들이 AI와 협력하면서도 자신의 판단력과 주도성을 유지할 수 있도록 학습 구조와 변화 촉진 환경을 마련하는 것이다.

1. 정규 학습에서 적시 학습으로의 전환

전통적인 정규 교육과정은 여전히 중요한 의미를 가진다. 조직의 공통 가치, 윤리 기준, 보안 원칙, 직무 기초 개념, 법적·제도적 준수 사항과 같은 내용은 여전히 체계적이고 계획된 교육 형식을 통해 전달될 필요가 있다. 특히 AI 기술이 빠르게 확산되는 환경에서는 오히려 이러한 기초 규범 교육의 중요성이 더 커질 수 있다. 왜냐하면 구성원들이 새로운 도구를 자유롭게 활용하는 상황일수록, 무엇이 허용되고 무엇이 위험한지에 대한 공통 기준이 더욱 중요해지기 때문이다. 그러나 이러한 정규 학습만으로는 실제 현장에서 발생하는 복잡하고 예측 불가능한 문제 상황에 충분히 대응하기 어렵다. 업무 환경의 변화 속도가 빠르고, 문제 유형이 세분화되며, 기술이 지속적으로 갱신되는 상황에서는 사전에 모든 내용을 가르쳐 두는 방식이 점점 비효율적이 된다. 이러한 이유로 HRD는 정규 학습의 가치를 부정하기보다, 그것만으로는 충분하지 않다는 점을 인식하고 적시 학습(just-in-time learning)의 비중을 대폭 확대할 필요가 있다(Brandenburg & Ellinger, 2003).

적시 학습의 핵심은 학습을 별도의 시간과 공간으로 떼어놓는 것이 아니라, 업무

흐름 내부에 삽입하는 데 있다. 구성원이 실제 문제에 직면했을 때 필요한 지식과 지원이 즉시 제공될수록, 학습은 추상적 지식 습득이 아니라 문제 해결과 연결된 살아 있는 경험이 된다. 이는 단지 편의성의 문제가 아니라 학습 전이의 문제이기도 하다. 전통적 집합교육은 교육 당시에는 이해가 이루어지더라도, 시간이 지나 실제 적용 맥락에 들어서면 상당 부분 망각되거나 전이되지 못하는 한계를 보여 왔다. 반면 적시 학습은 학습과 적용의 시간 간극을 최소화함으로써, 지식이 실제 행동으로 옮겨질 가능성을 높인다. 여기에 간격 반복(spaced repetition) 구조를 결합하면, 단발적 도움 제공을 넘어 장기 기억과 재인출을 촉진하는 학습 구조를 설계할 수 있다(Cepeda et al., 2008). 즉, HRD는 필요한 순간의 즉각적 지원과 이후의 복습·재강화를 하나의 연속된 체계로 엮어야 한다.

생성형 AI는 이러한 적시 학습을 구현하는 데 매우 강력한 기반이 될 수 있다. 그러나 여기서 중요한 것은 AI를 단순한 질의응답 도구나 요약 기계로 바라보지 않는 것이다. 만약 AI가 사용자의 질문에 대해 즉석 답변만 제공하는 수준에 머문다면, 그것은 편리한 검색 보조 도구일 수는 있어도 깊이 있는 학습을 촉진하는 장치가 되기는 어렵다. 반대로 AI가 사용자의 현재 상황을 해석하고, 문제 해결에 필요한 핵심 개념을 짧고 명료하게 설명하며, 예시와 체크리스트를 제공하고, 나아가 사용자가 스스로 판단해 보도록 반성 질문을 던지는 구조로 설계된다면, AI는 단순한 정보 제공자를 넘어 수행지원형 학습 장치로 기능할 수 있다(Molenaar, 2022). 이 경우 AI는 정답을 대신 제시하는 도구가 아니라, 사용자의 이해를 구조화하고 사고를 확장하는 실시간 코치에 가깝다.

이러한 이유로 HRD는 생성형 AI를 ‘언제든 물어보는 검색창’으로 도입해서는 안 되며, 오히려 마이크로러닝과 수행지원(performance support)측면을 통합하는 인프라로 설계해야 한다. 예컨대 업무 화면 내부에서 바로 확인 가능한 짧은 개념 카드, 상황별 체크리스트, 사례 기반 미니 가이드, 자주 발생하는 오류에 대한 경고 메시지, 사후 복습을 위한 자동 리마인더 등이 하나의 흐름으로 연결되어야 한다. 이렇게 되면 학습은 교육이 끝난 뒤 따로 떠올려야 하는 부가 활동이 아니라, 실제 업무 과정 속에서 자연스럽게 호출되고 반복되는 구조가 된다. 더 나아가 HRD전문가는 어떤 시점에 어떤 정보가 제공될 때 학습 효과가 가장 큰지, 사용자가 어떤 종류의 도움을 가장 자주 요청하는지, 도움을 받은 뒤 실제 수행이 개선되었는지를 지속적으로 분석하여 적시 학습 시스템 자체를 학습하는 구조로 고도화해야 한다. 결국 AI 시대 적시 학습의 핵심은 정보 접근 속도가 아니라, 업무

수행과 학습 전이가 하나의 경험으로 연결되도록 만드는 설계 능력에 있다.

2. 맞춤형 스킬 갭 진단 및 학습 생태계 구축

조직 차원의 역량 제고를 위해서는 더 이상 교육 수요를 사후적으로 파악하고 그에 따라 과정을 개설하는 방식만으로는 충분하지 않다. AI 시대에는 직무 요구가 훨씬 더 빠르게 바뀌고, 동일한 직무명 아래에서도 실제 과업 구성이 지속적으로 재편되기 때문이다. 이때 HRD가 담당해야 할 중요한 과제는 현재 조직이 가진 역량과 앞으로 요구될 역량 사이의 간극을 데이터 기반으로 조기에 식별하고, 그 결과를 실질적인 학습 개입으로 연결하는 것이다. 즉, 스킬 갭 진단은 단순한 진단 자체가 목적이 아니라, 미래의 일 변화에 대비한 선제적 학습 설계를 가능하게 하는 기반이 되어야 한다.

이와 관련한 선행 연구들은 스킬 데이터, 성과 데이터, 인사 데이터, 업무 데이터의 통합이 학습 설계의 정교화를 가능하게 한다고 지적한다(Cho et al., 2023; Wang et al., 2024). 예를 들어 어떤 부서에서 반복적으로 오류가 발생하는 지점, 특정 역할군에서 성과 편차가 크게 나타나는 패턴, 협업 과정에서 지연이 자주 발생하는 단계, 혹은 특정 업무 변화 이후 새로운 기술 활용 수준이 낮게 나타나는 현상은 모두 잠재적 학습 수요를 시사하는 신호가 될 수 있다. 문제는 이러한 데이터를 단순히 많이 모으는 것으로는 충분하지 않다는 점이다. 데이터의 품질이 낮거나, 지표가 실제 업무를 제대로 반영하지 못하거나, 결과가 구성원에게 설명되지 않으면 데이터 기반 진단은 신뢰를 얻기 어렵다. 더 나아가 AI 기반 추천이 어떤 논리에 의해 이루어졌는지 이해할 수 없다면, 구성원은 이를 지원이 아니라 감시나 분류의 장치로 받아들일 가능성이 높다. 따라서 맞춤형 학습 생태계의 핵심은 데이터 축적 그 자체가 아니라, 신뢰 가능한 해석과 실행 가능한 개입이다.

또한 AI 기반 스킬 갭 진단은 효율성의 문제만으로 접근해서는 안 된다. AI를 활용해 역량 부족 영역을 빠르게 식별할 수 있다는 사실은 분명 매력적이지만, 동시에 편향, 프라이버시, 설명 가능성, 인간 판단의 역할 같은 문제를 함께 고려해야 한다(Tambe et al., 2019). 예컨대 성과 데이터만을 기준으로 특정 직원군을 ‘고위험’ 또는 ‘저성과’ 집단으로 분류한다면, 그 결과는 낙인 효과를 낳을 수 있으며, 실제로는 구조적 자원 부족이나 역할 불명확성에서 비롯된 문제를 개인의 학습 부족으로 오인할 위험도 있다. 따라서 HRD는 스킬 갭 진단을 개인 선별이나

통제의 도구로 사용하기보다, 조직과 개인이 함께 성장하기 위한 학습 기회 설계의 출발점으로 활용해야 한다. 이는 곧 AI가 제시한 결과를 인간이 맥락적으로 해석하고, 당사자와 대화하며, 적절한 지원 방식으로 전환하는 절차가 반드시 필요함을 의미한다.

이러한 점에서 HRD는 스킬 갭 진단을 하나의 측정 프로젝트가 아니라 참여적 거버넌스가 내장된 학습 생태계 구축 프로젝트로 접근해야 한다. 여기서 중요한 질문은 단순히 “무엇이 부족한가”가 아니라, “어떤 데이터가 어떤 의사결정을 지원하는가”, “그 결과는 구성원에게 어떻게 설명되는가”, “그 설명은 어떤 학습 기회와 자원으로 환원되는가”이다. 예를 들어 특정 업무 변화에 따라 요구 역량이 달라졌다면, 그에 맞는 마이크로러닝, 코칭, 동료학습, 실습형 프로젝트, 현장 피드백이 어떻게 조합되어야 하는지를 함께 설계해야 한다. 이 과정에서 조직은 초기 파일럿 수준을 넘어, 진단-개입-성과 측정-재조정의 순환 구조를 갖추어야 하며, AI readiness와 프로세스 통합 수준도 함께 높여 나가야 한다(Johnk et al., 2021). 나아가 변화하는 일의 구조에 따라 HR 생태계 자체를 조정해야 한다는 점도 중요하다. HR 기능은 더 이상 채용, 평가, 교육이 분절된 형태로 작동해서는 안 되며, 변화하는 과업 구조와 학습 수요를 중심으로 보다 통합된 방식으로 재구성될 필요가 있다(Yalenios & d’ Armagnac, 2023).

3. 심리적 안전감이 담보된 실패 용인 환경 조성

AI 도입이 실제로 인간의 역량 증강으로 이어지기 위해서는 기술적 정교함 못지않게 심리적 안전감이 중요하다. 심리적 안전감은 팀 구성원이 질문하고, 반론을 제기하고, 오류를 보고하고, 새로운 방식을 실험하더라도 처벌이나 조롱을 당하지 않을 것이라는 공유된 믿음을 의미한다(Edmondson, 1999). 이는 단순히 분위기가 좋은 조직을 뜻하는 것이 아니라, 학습과 혁신이 가능한 조직의 핵심 조건을 뜻한다(이가현 & 임삭혁, 2018). 특히 AI 환경에서는 새로운 도구를 사용해보는 과정에서 오류, 혼란, 과잉 신뢰, 잘못된 위임, 데이터 입력 실수 등이 반복적으로 발생할 수 있기 때문에, 이러한 실수들을 숨기게 만드는 조직은 오히려 더 큰 위험을 축적하게 된다. 반대로 구성원들이 AI 활용 과정의 실패와 한계를 공개적으로 공유할 수 있을 때, 조직은 개별 실수를 집단적 학습 자원으로 전환할 수 있다.

이 점에서 AI 시대의 실패 용인 환경은 단순한 관용의 문제가 아니라, 책임 있는

실험을 가능하게 하는 제도적 조건의 문제다. 최근 연구가 보여주듯, 알고리즘 기반 업무 관리와 감시가 강화될수록 근로자는 자율성 저하, 스트레스 증가, 프라이버시 침해, 책임 전가의 위험에 직면할 수 있다(Todoli-Signes, 2021; Glavin et al., 2024). 만약 구성원들이 AI 도구 사용 과정의 모든 시행착오가 평가와 통제의 근거로 활용될 것이라고 느낀다면, 그들은 새로운 시도를 피하고, 안전한 선택만 하며, 결과적으로 조직은 AI 활용에 따른 학습 기회를 잃게 된다. 따라서 HRD는 AI 도입을 성과관리나 감시 강화와 단순 연결해서는 안 되며, 오히려 실험과 오류가 학습으로 환류될 수 있는 구조를 먼저 마련해야 한다.

이를 위해 HRD는 기술 조직과 협력하여 AI 활용의 원칙과 경계를 분명히 해야 한다. 첫째, 구성원은 시스템의 판단 근거를 일정 수준에서 설명받을 수 있어야 한다. 설명 가능성이 전혀 확보되지 않은 시스템은 도구가 아니라 블랙박스 경험되며, 이는 수용성과 신뢰를 동시에 저해한다. 둘째, 이의 제기와 인간 재검토가 가능한 절차가 마련되어야 한다. AI의 결과가 자동으로 최종 판단이 되어서는 안 되며, 특히 인사, 평가, 배치, 학습 추천처럼 개인에게 실질적 영향을 주는 사안에서는 인간의 검토와 수정 가능성이 제도적으로 보장되어야 한다. 셋째, 오류가 발생했을 때 안전하게 중단하고 인간이 통제권을 회수할 수 있는 인간중심 권한 구조가 필요하다. 이는 기술 실패에 대한 최소한의 안전장치일 뿐 아니라, 구성원들이 도구를 맹목적으로 따르지 않고 비판적으로 활용하도록 돕는 심리적 기반이 된다(Parker & Grote, 2022; Glikson & Woolley, 2020). 넷째, 무엇보다 중요한 것은 실험과 실패를 학습으로 전환하는 피드백 문화다. 실패를 허용한다는 것은 기준 없이 방임한다는 뜻이 아니다. 오히려 어떤 실패는 허용 가능하고 어떤 실패는 치명적인지, 실패가 발생했을 때 무엇을 기록하고 누구와 공유하며 어떻게 개선할 것인지를 명확히 구조화해야 한다. 예를 들어 AI 활용 중 발생한 환각 사례, 부적절한 자동화 위임 사례, 편향된 추천 사례 등을 개인의 실수로만 처리하지 않고, 조직 차원의 사례 저장소와 리뷰 세션, 사후 성찰 문서, 개선 가이드라인으로 연결하는 절차가 필요하다. 이렇게 될 때 실패는 숨겨야 할 오점이 아니라, 다음 설계를 더 정교하게 만드는 학습 자원이 된다.

V. 결론

AI 시대의 일은 단순히 인간의 업무 일부를 기계에 이전하는 과정이 아니다.

오히려 인간이 더 높은 수준의 판단, 검증, 조율, 학습을 수행할 수 있도록 일의 구조 자체를 재설계하는 과정에 가깝다. 생성형 AI와 자동화 기술은 분명 반복적이고 표준화된 과업의 부담을 줄여 주지만, 동시에 결과의 타당성을 확인하고, 복잡한 맥락 속에서 의미를 해석하며, 윤리적·전략적 책임을 지는 인간의 역할을 더욱 선명하게 부각시킨다. 따라서 미래의 경쟁력은 단순히 더 많은 AI 도구를 보유하는 데서 나오지 않으며, 인간과 기술의 기여를 얼마나 정교하게 결합하고 조율할 수 있는가에 의해 결정된다.

이러한 변화 속에서 HRD의 역할 역시 근본적으로 확장되어야 한다. HRD는 더 이상 교육 프로그램을 기획·운영하는 지원 기능에 머물러서는 안 되며, 인간-AI 협업이 효과적으로 이루어질 수 있도록 직무를 재설계하고, 조직 구성원이 기술을 비판적으로 활용할 수 있도록 학습 환경을 조성하며, 학습자 에이전시와 심리적 안전감을 보호하는 사회기술적 설계자로 자리매김해야 한다. 또한 데이터 거버넌스, 설명 가능성, 책임 있는 활용 원칙을 포함한 조직 차원의 제도적 기반을 함께 구축함으로써, 기술 활용이 단기적 효율성에 머무르지 않고 지속 가능한 학습과 성장으로 이어지도록 만들어야 한다.

결국 AI 시대 HRD의 핵심 과제는 기술을 빠르게 도입하는 것 자체가 아니라, 그 기술이 인간의 사고와 학습을 약화시키지 않고 오히려 증폭시키도록 만드는 견고한 사회기술적 토대를 설계하는 데 있다. 다시 말해, AI 시대 HRD의 본질은 기술 적응을 지원하는 데 그치지 않고, 인간이 기술과의 협업 속에서도 판단의 주체이자 학습의 주체로 남을 수 있도록 조직의 조건을 설계하는 데 있다. 바로 그 지점에서 HRD는 단순한 교육 부서를 넘어, AI 시대 조직의 지속 가능성과 인간 중심 혁신을 떠받치는 핵심 전략 기능으로 재정의될 필요가 있다.

참고문헌

- 이가현, 임상혁. (2018). 인사담당자의 심리적 안전감이 혁신행동에 미치는 영향: 긍정심리자본의 조절효과: 긍정심리자본의 조절효과. *대한경영학회지*, 31(11), 2125-2145.
- Abhari, K. (2025). Employee participation in digital transformation: From digitalization sentiment to transformation predisposition. *Information & Management*, 62(8), 104212. <https://doi.org/10.1016/j.im.2025.104212>
- Acemoglu, D., & Restrepo, P. (2019). Automation and new tasks: How technology displaces and reinstates labor. *Journal of Economic Perspectives*, 33(2), 3-30. <https://doi.org/10.1257/jep.33.2.3>
- Autor, D. H., & Handel, M. J. (2013). Putting tasks to the test: Human capital, job tasks, and wages. *Journal of Labor Economics*, 31(S1), S59-S96. <https://doi.org/10.1086/669332>
- Autor, D. H., Levy, F., & Murnane, R. J. (2003). The skill content of recent technological change: An empirical exploration. *The Quarterly Journal of Economics*, 118(4), 1279-1333. <https://doi.org/10.1162/003355303322552801>
- Brandenburg, D. C., & Ellinger, A. D. (2003). The future: Just-in-time learning expectations and potential implications for human resource development. *Advances in Developing Human Resources*, 5(3), 308-320. <https://doi.org/10.1177/1523422303254629>
- Brynjolfsson, E., Li, D., & Raymond, L. (2025). Generative AI at work. *The Quarterly Journal of Economics*, 140(2), 889-942. <https://doi.org/10.1093/qje/qjae044>
- Cepeda, N. J., Vul, E., Rohrer, D., Wixted, J. T., & Pashler, H. (2008). Spacing effects in learning: A temporal ridge of optimal retention. *Psychological Science*, 19(11), 1095-1102. <https://doi.org/10.1111/j.1467-9280.2008.02209.x>
- Cho, W., Choi, S., & Choi, H. (2023). Human resources analytics for public personnel management: Concepts, cases, and caveats. *Administrative Sciences*, 13(2), Article 41. <https://doi.org/10.3390/admsci13020041>
- Dellermann, D., Ebel, P., Sollner, M., & Leimeister, J. M. (2019). Hybrid intelligence. *Business & Information Systems Engineering*, 61(5), 637-643.

- <https://doi.org/10.1007/s12599-019-00595-2>
- Edmondson, A. (1999). Psychological safety and learning behavior in work teams. *Administrative Science Quarterly*, 44(2), 350–383.
<https://doi.org/10.2307/2666999>
- Fugener, A., Grahl, J., Gupta, A., & Ketter, W. (2022). Cognitive challenges in human-artificial intelligence collaboration: Investigating the path toward productive delegation. *Information Systems Research*, 33(2), 678–696.
<https://doi.org/10.1287/isre.2021.1079>
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660.
<https://doi.org/10.5465/annals.2018.0057>
- Georganta, E., & Ulfert, A. S. (2024). Would you trust an AI team member? Team trust in human-AI teams. *Journal of Occupational and Organizational Psychology*, 97(3), 1212–1241. <https://doi.org/10.1111/joop.12504>
- Gillath, O., Ai, T., Branicky, M. S., Keshmiri, S., Davison, R. B., & Spaulding, R. (2021). Attachment and trust in artificial intelligence. *Computers in Human Behavior*, 115, Article 106607. <https://doi.org/10.1016/j.chb.2020.106607>
- Glavin, P., Bierman, A., & Schieman, S. (2024). Private eyes, they see your every move: Workplace surveillance and worker well-being. *Social Currents*, 11(4), 327–345. <https://doi.org/10.1177/23294965241228874>
- GPAI (2025). Generative AI and the future of work global dialogue: Perceptions and prospects. Roundtables in Asia, Europe, and Latin America, Retrieved April 10, 2026, from https://wp.oecd.ai/app/uploads/2025/01/20250128_GPAI_GenAI_FoW_report_final_VOECD.pdf
- Holmstrom, J. (2022). From AI to digital transformation: The AI readiness framework. *Business Horizons*, 65(3), 329–339. <https://doi.org/10.1016/j.bushor.2021.03.006>
- Ivanov, D., & Dolgui, A. (2021). A digital supply chain twin for managing the disruption risks and resilience in the era of Industry 4.0. *Production Planning & Control*, 32(9), 775–788. <https://doi.org/10.1080/09537287.2020.1768450>
- Jarrahi, M. H. (2018). Artificial intelligence and the future of work: Human-AI

- symbiosis in organizational decision making. *Business Horizons*, 61(4), 577–586.
<https://doi.org/10.1016/j.bushor.2018.03.007>
- Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y., Chen, D., Dai, W., Chan, H. S., Madotto, A., & Fung, P. (2023). Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), Article 248.
<https://doi.org/10.1145/3571730>
- Johnk, J., Weißert, M., & Wyrski, K. (2021). Ready or not, AI comes—An interview study of organizational AI readiness factors. *Business & Information Systems Engineering*, 63(1), 5–20. <https://doi.org/10.1007/s12599-020-00676-7>
- Klingbeil, A., Grutzner, C., & Schreck, P. (2024). Trust and reliance on AI—An experimental study on the extent and costs of overreliance on AI. *Computers in Human Behavior*, 160, 108352. <https://doi.org/10.1016/j.chb.2024.108352>
- Li, J., & Yeo, R. K. (2024). Artificial intelligence and human integration: A conceptual exploration of its influence on work processes and workplace learning. *Human Resource Development International*, 27(3), 367–387.
<https://doi.org/10.1080/13678868.2024.2348987>
- Mikalef, P., & Gupta, M. (2021). Artificial intelligence capability: Conceptualization, measurement calibration, and empirical study on its impact on organizational creativity and firm performance. *Information & Management*, 58(3), Article 103434. <https://doi.org/10.1016/j.im.2021.103434>
- Molenaar, I. (2022). Towards hybrid human-AI learning technologies. *European Journal of Education*, 57(4), 632–645. <https://doi.org/10.1111/ejed.12527>
- Noy, S., & Zhang, W. (2023). Experimental evidence on the productivity effects of generative artificial intelligence. *Science*, 381(6654), 187–192.
<https://doi.org/10.1126/science.adh2586>
- Orso, V., Ziviani, R., Bacchiega, G., Bondani, G., Spagnolli, A., & Gamberini, L. (2022). Employee-centric innovation: Integrating participatory design and video-analysis to foster the transition to Industry 5.0. *Computers & Industrial Engineering*, 173, 108661. <https://doi.org/10.1016/j.cie.2022.108661>
- Overby, E., Bharadwaj, A., & Sambamurthy, V. (2006). Enterprise agility and the enabling role of information technology. *European Journal of Information*

- Systems*, 15(2), 120–131. <https://doi.org/10.1057/palgrave.ejis.3000600>
- Parker, S. K., & Grote, G. (2022). Automation, algorithms, and beyond: Why work design matters more than ever in a digital world. *Applied Psychology*, 71(4), 1171–1204. <https://doi.org/10.1111/apps.12241>
- Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation–augmentation paradox. *Academy of Management Review*, 46(1), 192–210. <https://doi.org/10.5465/amr.2018.0072>
- Seeber, I., Bittner, E., Briggs, R. O., de Vreede, T., de Vreede, G. J., Elkins, A., Maier, R., Merz, A. B., Oest-Rei ß, S., Randrup, N., Schwabe, G., & Sollner, M. (2020). Machines as teammates: A research agenda on AI in team collaboration. *Information & Management*, 57(2), Article 103174. <https://doi.org/10.1016/j.im.2019.103174>
- Teece, D. J., Pisano, G., & Shuen, A. (1997). Dynamic capabilities and strategic management. *Strategic Management Journal*, 18(7), 509–533
- Tambe, P., Cappelli, P., & Yakubovich, V. (2019). Artificial intelligence in human resources management: Challenges and a path forward. *California Management Review*, 61(4), 15–42. <https://doi.org/10.1177/0008125619867910>
- Todoli-Signes, A. (2021). Making algorithms safe for workers: Occupational risks associated with work managed by artificial intelligence. Transfer: *European Review of Labour and Research*, 27(4), 433–452. <https://doi.org/10.1177/10242589211035040>
- Trocin, C., Hovland, I. V., Mikalef, P., & Dremel, C. (2021). How artificial intelligence affords digital innovation: A cross-case analysis of Scandinavian companies. *Technological Forecasting and Social Change*, 173, Article 121081. <https://doi.org/10.1016/j.techfore.2021.121081>
- Ullrich, A., Rei ßig, M., Niehoff, S., & Beier, G. (2023). Employee involvement and participation in digital transformation: A combined analysis of literature and practitioners' expertise. *Journal of Organizational Change Management*, 36(8), 29–48. <https://doi.org/10.1108/JOCM-10-2022-0302>
- Wang, L., Zhou, Y., Sanders, K., Marler, J. H., & Zou, Y. (2024). Determinants of effective HR analytics implementation: An in-depth review and a dynamic

framework for future research. *Journal of Business Research*, 170, Article 114312. <https://doi.org/10.1016/j.jbusres.2023.114312>

Yalenios, J., & d' Armagnac, S. (2023). Work transformation and the HR ecosystem dynamics: A longitudinal case study of HRM disruption in the era of the 4th industrial revolution. *Human Resource Management*, 62(1), 55–77. <https://doi.org/10.1002/hrm.22114>

김동호 교수

現 서울대학교 산업인력개발학과 교수

前 성균관대학교 교수

前 플로리다대학교 교수

前 노던일리노이대학교 교수

관심분야: AI시대의 일, 학습, 경력개발,
다중 에이전틱 AI 시스템, 퍼플애널리틱스





발행인 편집인 : 조대연
인쇄 : 고려대학교 HRD 정책연구소
이메일 : Kuhrd@korea.ac.kr