

한국 SCM 학회지

*Journal of the Korean Society of
Supply Chain Management*

Volume 21 Number 3

2021 December

디지털 공급망의 이론 및 응용 특집호

 사단
법인 한국SCM학회

한국 SCM 학회지

Journal of the Korean Society of Supply Chain Management

디지털 공급망의 이론 및 응용

1. 총완료시간 최소화를 위한 3-설비 중첩 대기시간 제약 플로우샵 스케줄링

이준호 · 김현정

We examine a three-machine flow shop scheduling problem in which two overlapping waiting time limits exist between the first and second machines and the first and third machines, respectively. The objective is to minimize the total completion time. We first present a mixed integer linear programming (MILP) to mathematically model the problem and obtain optimal solutions. Due to the high computational complexity of the problem, large-sized problems cannot be solved by the MILP within a reasonable computation time. We therefore develop a heuristic scheduling method based on the genetic algorithm. Finally, the performance of the proposed algorithm is numerically tested.

13. 빅데이터 DEA를 위한 기계학습의 적용

응위엔투이즈영 · 임성복

Data envelopment analysis (DEA) is a tool for identifying best-practices when multiple performance metrics or measures are present for decision-making units (DMUs). As big data issue becomes an important area of supply chain and operations management, DEA is evolving into a data-oriented data science tool for benchmarking, performance evaluation, composite index construction and others. As the number of DMUs increases, the running-time to solve the standard DEA model sharply rises. Such situations are appearing more frequently in the era of big data. This issue could be an important challenge particularly when real-time data stream in at extremely high rates and the DEA analysis needs to be performed very quickly.

Therefore, there exist practical needs for developing an efficient way of solving large-scale DEA problems. In this paper, we propose a practical approach for speeding up the DEA efficiency estimation process based on machine learning. In this approach, a sample of DMUs is selected from the population as a training data set, based on which a machine is trained to predict the efficiency scores of unselected or newly streamed-in DMUs. We also suggest a data augmentation technique to enhance the learning process under severe data class imbalance. The superior performance of the proposed approach over the conventional one in terms of efficiency prediction power as well as model computation time is shown through a series of computational experiments using randomly generated data.

27. Industry 4.0 기술과 디지털 SCM 실행방식의 연계: 중국 시장을 중심으로

김영길 · 박정수

This study aims to check whether Digital Supply Chain Management practices has positive effect on company's performance by using sample consisting of companies in Chinese industry. Furthermore, the authors check if companies' utilizing either AI related technologies or Big Data analytics have positive moderating or accelerating effect on the relationship between DSCM and the performance. To achieve these goals, the researchers conducted surveys and empirical researches using validity check, reliability analysis, regression model and moderate regression model. Results of further research showed that DSCM practices affects the performance positively in Chinese sample companies. As the further research results, utilizing AI related technologies has positive moderating effect on relationship between DSCM practices and the performance, while utilizing Big Data analytics does in limited range. From these research results a managerial implication that implementing DSCM

practices and AI related technologies at the same time gives Chinese companies more improvement in performances due to the synergistic effects between them

37. 혼성제조시스템을 위한 생산용량할당정책 수립과 지연재고정책의 이익 성과 효과성 분석

김은갑

We consider a hybrid manufacturing system that produces components based on both long-term contract with higher priority and short-term contract with lower priority. The manufacturer operates production on a make-to-stock mode for the long-term contract and on a make-to-order mode for the short-term contract. Recently, hybrid manufacturing has received attention from academia and industry because it can contribute to profit enhancement through revenue channel diversification. In this paper, using a Markov decision process model, we study the structure of capacity rationing and admission control policies under the assumption of non-preemptive batch production and compound Poisson demand process. We also investigate the effectiveness of backlogging short-term contract orders on profit and examine the impact of long-term contract demand sizes on profit with varying probabilities of each demand size and variance of demand.

45. SNA를 이용한 제조용 로봇산업 디지털 공급망 가치사슬 분석

김태우 · 서창교

Industrial robots are manufactured as robot products by utilizing parts including sensors, actuators, controllers and consist of a supply

chain provided for end-users through robotics system integrator. This study was analyzed with a social network analysis(SNA) by using corporate trade data to find the characteristics of the industrial robot supply chain in Korea. The network of the industrial robot supply chain consist of 3200 nodes and 3488 edges and the key types of demand industry have found to be automobiles, electrical and electronic equipment, and metal and machinery. The characteristics of the supply chain of industrial robot industry in Korea were analyzed that foreigner-invested robot companies are very influential. The size of sales and importance of the network do not accord with each other. It is analyzed that the companies are classified into the ones related to robots, but their sales in the robot sector is remarkably low or there are transactions with the ones which have nothing to do with them. In order to improve the Korean industrial robot industry's competitiveness, the implementation of the systems for specialized in robot fields and systematic management of information on robot-related companies are required and the existing quantitative polices for supplying and diffusion industrial robots need to be implemented with qualitative polices for utilizing and promoting them.

61. 재생에너지 생산시스템에서 검사 자동화의 효과 연구

Mitali Sarkar · 서용원

With the changing lifestyle of humans, the use of technologies increases, thus increasing the necessity of energy demand. Owing to limited non-renewable energy sources, the demand for renewable energy is growing gradually. On the other hand, the waste generated increases with the increasing population rate. Wastes can be a useful source of energy demand, and environmental pollution can be reduced. This study explains a renewable energy production system. The proposed production system uses wastes as

a source of raw material. Wastes are collected and transported from the collection center of a city to produce renewable energy. An automated inspection policy is used in the production system. The random rate of unusable wastes is scraped from the system, and the rest are sent for energy production. The total cost of the energy conversion plant is optimized analytically by a classical optimization method. Three numerical examples are given for uniform, triangular, and beta distribution of random scrap rate. The results show that adopting beta distribution for the proposed study is required to obtain the maximum amount of renewable energy at the minimum cost. The transportation cost is found as the most sensitive cost parameter among all cost parameters. The proposed study fulfills both the aims of renewable energy production and waste management.

71. 데이터기반 재고관리를 위한 DDPG기반 심층강화학습 알고리즘 연구

이병권 · 박건수 · 정세윤

We develop a modified version of deep deterministic policy gradient(DDPG) algorithm, which is one of the most popular deep reinforcement learning algorithms dealing with continuous action spaces, and apply it to the infinite-horizon newsvendor problems with constant lead time. Reflecting the key features of the inventory management problems, our algorithm, named as Inventory based DDPG (IDDPG), is differentiated in the state variables, action spaces, and cost functions compared to the DDPG. By conducting numerical experiments and comparing the performances, we found that when applied to the inventory management problems, IDDPG is able to find the near-optimal solutions and outperforms the DDPG in most cases. We also found that our IDDPG can be applied to a inventory management problem whose optimal solution is not known.

한국 SCM 학회지

Journal of the Korean Society of Supply Chain Management

투고논문 작성 요령

1. 제출방법

투고자는 논문을 한글 또는 MS워드 작성하며, 글씨크기 11포인트, 2단 편집으로 작성하여 제출한다. 논문 심사 후 게재가 확정되면 저자 약력 및 사진이 포함된 최종본을 e-mail로 제출하여야 한다(논문저자 중 한 명 이상은 한국 SCM학회 연회비 납부회원이어야 투고할 수 있다).

2. 제출 철회

접수된 후 심사과정에 있는 논문의 철회를 저자가 원하는 경우, 저자는 서면으로 편집위원장에게 철회요청서를 제출하여야 한다.

3. 표지 및 본문의 내용

논문 표지에는 국문으로 논문 제목, 저자명, 소속을 기입하고 영문으로 다시 논문 제목, 저자명, 소속을 기입한 후 영문 요약, 키워드를 차례로 기입한다. 표지 각주에는 사사표기, 교신저자 정보(영문 주소, 전화번호, e-mail 주소)를 기입한다. 표지의 다음 쪽부터 본문, 참고문헌, 부록 순으로 작성한다. 원고 작성 시 본문과 그림(그래프) 등은 모두 흑백으로 작성한다(컬러 그림(그래프) 사용 자제).

4. 영문 작성

영문의 대문자는 고유명사나 문장의 첫 자 또는 고유명사의 약자 등에만 사용한다.

5. Abstract 및 키워드

영문으로 기입된 저자 소속 아래 150단어 이내의 영문 요약(Abstract)을 기입하고, 그 아래 키워드를 5개 내외로 정해 영문으로 기입한다.

6. 각주(footnote)

- 연구비의 지원을 받아 연구가 이루어진 논문을 알릴 경우
- 교신저자의 연락처를 기재하는 경우
상기 2개 사항을 제외하고, 각주는 사용하지 않는 것을 원칙으로 한다.

7. 저자 구분

논문의 저자 기재 시 제1저자, 제2저자 순으로 기재하며, 교신저자의 경우 “+” 표시를 이름 옆에 표기한다.

8. 본문 제목 일련번호 표기방법

장, 절, 항은 아라비아 숫자로 1., 1.1., 1.1.1., 등으로 표기한다.

9. 그림과 표

그림과 표는 제목과 내용을 모두 영문으로 작성한다. 그림 제목은 “Fig.”로 표기한 후 일련번호를 매기고, 그림 아래 가운데 정렬한다. 표의 제목은 “Table”로 표기한 후 일련번호를 매기고, 표 위 왼쪽 정렬한다. 모든 그림과 표는 본문의 적당한 위치에 삽입하고, 삽입이 어려운 경우에는 논문의 맨 뒤에 첨부한다.

10. 수식 표현

수식은 필요한 경우 아래 예시와 같이 오른쪽에 (1), (2), (3), ... 등으로 일련번호를 부여해 표기한다.

예시) $y = a_1x^2 + a_2x + a_3$ (1)

11. 참고문헌

참고문헌의 표제는 “REFERENCES”로 표기하고, 이하의 모든 내용은 영문으로 작성한다. 참고문헌은 [1], [2], ... 등의 일련번호를 붙여 알파벳순으로 나열한다. 인용된 문헌의 종류별 세부 작성 방식은 다음과 같다.

예시)

- 학술지

[1] Hayes, R. & Pisano, G. P.(2000). SCM Strategy in Korea. *SCM Journal*, 11(4), 25-41.

- 단행본

[2] Hayes, R.(2000). *SCM Strategy in Korea* (2nd ed). New Jersey: Prentice-Hall.

12. 논문 심사료 및 게재료

심사료는 5만원, 게재료는 10페이지(2단으로 편집된 최종 게재본 기준)를 기본으로 20만원이고, 10페이지 초과 시 페이지 당 2만원을 추가로 납부한다. 또한 각주 중 연구비 지원에 대한 사사표기가 있을 경우에는 10만원을 추가로 납부한다(연구비 지원 금액이 1천만 원을 넘지 않을 경우 사무국에 사사표기 금액 면제 요청 가능).

<송금처>

국민은행 031737-00-000482

(예금주: 사단법인 한국SCM학회 / 영수증 발급)

디지털 공급망의 이론 및 응용

존경하는 한국SCM학회 회원 여러분, 안녕하십니까?

최근 고객 수요의 개인화, 다양화, 비대면화를 위한 서비스 요구와 함께 온라인 쇼핑 등 다수의 구매 채널을 동시에 관리하기 위한 효과적인 수단으로 공급망의 디지털화에 대한 관심이 증가하고 있습니다. 공급망의 디지털화는 공급자와 물류 시스템에서 발생하는 각종 위험에 대비하고, 최근 발전하고 있는 빅데이터, IoT, 인공지능 등의 변혁적 기술(disruptive technology)들을 적극적으로 활용하여 공급망의 효율성을 높이는 데 활용되고 있습니다.

본 특집호는 새롭게 부상하는 디지털 공급망의 개념 정립, 디지털 공급망 관련 기초 및 응용 기술, 산업 및 정책 사례 등에 관한 최신의 연구 결과를 정리하였습니다. 특집호에 투고된 논문들의 심사를 마친 결과 총 7편의 논문을 게재하게 되었습니다. 본 자리를 빌어 소중한 연구 논문을 보내주신 저자분들과 바쁜 와중에도 성실하게 논문을 검토해 주신 심사자분들께 깊은 감사를 전합니다.

모쪼록 이번 특집호가 한국SCM학회지의 발전과 디지털 공급망 연구의 확산에 도움이 되기를 바라며 본 특집호의 출간사를 마칩니다. 감사합니다.

한국SCM학회지 편집위원장

동국대학교 경영대학 임성목

서울대학교 산업공학과 박진수

총완료시간 최소화를 위한 3-설비 중첩 대기시간 제약 플로우샵 스케줄링*

이준호** · 김현정***†

충남대학교 경영학부, *카이스트 산업및시스템공학과

3-Machine Flow Shop Scheduling with Overlapping Waiting Time Constraints to Minimize Total Completion Time

Jun-Ho Lee** · Hyun-Jung Kim***†

School of Business, Chungnam National University, *Department of Industrial & Systems Engineering, KAIST

We examine a three-machine flow shop scheduling problem in which two overlapping waiting time limits exist between the first and second machines and the first and third machines, respectively. The objective is to minimize the total completion time. We first present a mixed integer linear programming (MILP) to mathematically model the problem and obtain optimal solutions. Due to the high computational complexity of the problem, large-sized problems cannot be solved by the MILP within a reasonable computation time. We therefore develop a heuristic scheduling method based on the genetic algorithm. Finally, the performance of the proposed algorithm is numerically tested.

Keyword : Scheduling, Genetic Algorithm, Flow Shop, Waiting Time Constraint

* 본 연구는 충남대학교 학술연구비에 의해 지원되었음.

† **Corresponding Author** : Department of Industrial & Systems Engineering, KAIST, 291 Daehak-ro, Yuseong-Gu, Daejeon 34141, Korea, Tel: +82-42-350-3134,
E-mail: hyunjungkim@kaist.ac.kr,

Received: 22 October 2021, Accepted: 7 December 2021

1. 서 론

플로우샵은 생산 및 제조에서 가장 널리 사용되는 시스템 중 하나이며, 그 중요성으로 인해 지금까지 많은 스케줄링 연구가 수행되어져 왔다(An et al., 2016; Kim & Lee, 2019; Rajendran & Chaudhuri, 1992; Aldowaisan & Allahverdi, 2004; Lee, 2020; Yu et al., 2017; Lee et al., 2021). 플로우샵 스케줄링 연구들은 대부분 특정 목적식(objective)을 최소화(또는 최대화)하기 위해 플로우샵을 구성하는 각 설비에서 작업의 공정순서와 공정시작 시각을 결정하는 문제를 다룬다. 플로우샵 스케줄링 문제에서 주로 사용되는 목적식으로는 최대완료시간(makespan) 최소화, 총완료시간(total completion time; TCT) 최소화, 총가중완료시간(total weighted completion time) 최소화, 납기 지연(tardiness) 최소화 등이 있다(Pinedo, 2016; Kim & Lee, 2019; Hong & Yoon, 2021). 본 논문에서는 목적식으로 총완료시간 최소화를 고려하며 기존 연구에서 많이 다루어지지 않았지만 실제 생산 환경에서 중요한 스케줄링 요구사항 중 하나인 대기시간 제약을 고려한다. 구체적으로 본 논문에서는 3-설비 플로우샵 스케줄링 문제를 다루며, 첫 번째 설비와 두 번째 설비 사이 그리고 첫 번째 설비와 세 번째 설비 사이에 각각 대기시간 제약 a_i 와 b_i 가 존재하는 중첩 대기시간 제약을 고려한다. a_i 와 b_i 에서 i 는 작업(job)을 나타낸다. 즉, 각각의 작업 별로 준수해야 하는 대기시간 제약이 상이한 일반적인 경우를 가정한다. 총 n 개의 작업이 있을 때 작업 $i \in \{1, 2, \dots, n\}$ 는 다음의 두 제약을 만족하도록 스케줄링이 되어야 한다: 1) 첫 번째 설비에서 공정이 완료된 후 두 번째 설비에 투입되기까지 대기시간이 a_i 이하여야 함, 2) 첫 번째 설비에서 공정이 끝난 후 세 번째 설비에 투입되어 공정이 시작되기 전까지 대기시간이 b_i 이내여야 함. 이러한 대기시간 제약은 작업물이 외부환경에 일정시간 이상 노출될 경우 품질저하가 발생할 수 있는 산업에서 빈번히 고려되며 대표적인 산업으로는 반도체, 배터리, 철강, 유제품 제조업 등이 있다(Kim & Lee, 2019; Lee et al., 2021; Ju et al., 2017; Wang et al., 2010; Lee et al., 2018). 또한 a_i 와 b_i 가 모두 0인 경우 no-wait 플로우샵 스케줄링 문제가 되며, 모두 충분히 큰 경우 대기시간 제약이 없는 일반적인 플로우샵 스케줄링 문제와 동일하다는 점에서 본 연구에서 다루는 문제가 매우 포괄적이라고 할 수 있다.

본 논문에서는 전술한 제약사항들을 고려하여 총완료시간을 최소화할 수 있는 혼합정수계획모형(mixed integer linear programming; MILP)을 제안한다. 본 논문에서 다루고자 하는 문제는 NP-complete이므로, 제안된 MILP를 통해 최적해를 도출할 수 있는 문제 크기에는 한계가 있다(Kim & Lee, 2019). 따라서 문제의 크기가 큰 경우에도 비교적 빠른 시간 내에 우수한 해를 도출할 수 있는 휴리스틱 알고리즘 또한 제안한다. 제안하는 휴리스틱 알고리즘은 유전알고리즘(genetic algorithm)을 기반으로 하며 성능을 높이기 위해 지역탐색(local search) 방법 또한 포함한다. 추가적으로 유전알고리즘의 경우 초기모집단(initial population)이 우수해야 결과적으로 좋은 해를 얻을 확률이 높아지기 때문에 우수한 초기모집단을 얻기 위해 기존 연구들을 참고하여 다양한 휴리스틱 방법들을 사용한다.

본 논문의 구성은 다음과 같다. 우선 2장에서 본 연구와 관련된 기존 연구들을 소개하고, 3장에서 본 연구에서 다루는 문제정의와 MILP를 제시한다. 4장에서는 유전알고리즘 기반의 휴리스틱 알고리즘을 제시하며, 5장에서는 실험결과를 소개한다. 마지막으로 6장에서는 결론을 제시하고 향후 연구 방향에 관해 논의한다.

2. 관련문헌

본 논문에서 다루고 있는 문제와 관련된 문헌을 크게 두 가지 방향으로 나누어 살펴본다. 첫 번째는 총완료시간 최소화를 위한 플로우샵 스케줄링에 관한 연구들이며, 두 번째는 본 연구에서 고려하는 중요한 스케줄링 요구사항인 중첩 대기시간 제약이 있는 플로우샵 스케줄링 문제에 관한 연구들이다.

플로우샵 스케줄링을 다룬 연구는 다수 존재한다(Pinedo, 2016). 합리적인 시간 안에 플로우샵 스케줄링 문제의 최적해를 도출하기 위해 Branch & Bound(B&B) 알고리즘 등을 개발한 연구도 존재하지만(An et al., 2016; Kim & Lee, 2019), 앞서 1장에서 언급한 바와 같이 문제의 복잡도가 높기 때문에 대부분의 연구들은 효율적인 휴리스틱 알고리즘 개발에 초점을 맞추었다.

목적식을 총완료시간 최소화로 국한해도 지금까지 다수의 효율적인 휴리스틱 알고리즘들이 개발되었다. Rajendran & Chaudhuri(1992)는 플로우샵에서 총완료시간 최소화를 위해 세 가지 휴리스틱 알고리즘을 제안하였으며, 기존 알고리즘보다 제안된 알고리즘이 우수한 성능을 보임을 실험적으로 규명하였다. Aldowaisan &

Allahverdi(2004)는 본 연구에서 다루는 문제의 특수한 경우라고 할 수 있는 no-wait 플로우샵 스케줄링 문제에 서 총완료시간 최소화를 위한 휴리스틱 알고리즘을 제안 하였다. 구체적으로 초기해를 얻기 위해 initial sequence algorithm(ISA)을 제안하였고, ISA를 바탕으로 해를 개선 시키는 추가 알고리즘들을 제안하였다. 동일한 문제에 대해 Gao et al.(2013)은 작업물이 각각의 설비에서 가지는 공정시간의 표준편차가 총완료시간에 영향을 주는 핵심 요소가 될 수 있음에 착안하여 improved standard deviation heuristic(ISDH)을 제안하였다. Laha & Sarin (2009)은 Framina & Leisten(2003)의 알고리즘을 개선하여 순열 플로우샵(permutation flow shop)에서 총완료시간 최소화를 위한 알고리즘을 제안하였다. 실험을 통하여 제안한 알고리즘이 기존에 우수하다고 알려졌던 Framina & Leisten(2003)의 알고리즘 보다 우수한 성능을 보임을 규명하였다. 앞서 소개한 효율적인 휴리스틱 알고리즘들은 본 연구에서 제안하는 유전알고리즘 기반 휴리스틱의 성능을 높이기 위한 초기모집단 구성에 활용 된다.

본 연구의 중요한 스케줄링 요구사항인 대기시간 제약을 플로우샵 스케줄링 문제에서 다룬 연구는 비교적 많지 않다. Joo & Kim(2009)는 대기시간 제약이 있는 2-설비 플로우샵의 makespan 최소화를 위한 B&B 알고리즘을 제안하였다. 동일한 문제에 대해 An et al.(2016)은 sequence-dependent 셋업시간을 추가로 고려한 B&B 알고리즘을 개발하였다. Lee(2020)는 대기시간 제약이 있는 2-설비 플로우샵에서 총납기 지연(total tardiness)을 최소화하기 위한 유전알고리즘을 개발하였고, Yu et al. (2017)은 유사한 환경에서 대기시간의 산포를 줄이기 위한 연구를 수행하였다. 전술한 연구들이 대기시간 제약을 고려하긴 하지만 본 연구에서 다루는 중첩 대기시간 제약과는 차이가 있으며 해당 결과를 본 연구에 적용하기에는 어려움이 있다.

중첩 대기시간 제약을 다룬 연구로는 Kim & Lee (2019)와 Lee et al.(2021)이 있다. Kim & Lee(2019)는 본 연구에서 고려하고 있는 중첩 대기시간 제약 하에서 makespan을 최소화하는 최적해 도출을 위한 B&B 알고리즘을 개발하였다. Lee et al.(2021)은 동일한 문제에 대하여 빠른 시간 안에 우수한 해를 도출할 수 있는 휴리스틱 알고리즘을 개발하였다.

전술한 바와 같이 플로우샵 스케줄링에 관한 연구는 다수 존재하지만, 총완료시간 최소화를 목적식으로 고려하며 중첩 대기시간 제약을 갖는 플로우샵 스케줄링 문제는 지금까지 다루어지지 않았다.

3. 문제정의 및 MLP

본 연구에서는 총완료시간 최소화를 위한 3-설비 중첩 대기시간 제약 플로우샵 스케줄링 문제를 다룬다. 3개의 설비 M1, M2, M3와 n 개의 작업 $1, 2, \dots, n$ 에 대하여 각각의 설비에서의 작업 순서가 동일한 순열 플로우샵을 가정한다. 예를 들어, 순열 플로우샵에서는 M1에서 $1, 2, \dots, n$ 의 순서로 작업이 이루어졌다면 M2와 M3에서도 $1, 2, \dots, n$ 의 순서로 작업이 이루어진다. 순열 스케줄이 항상 최적해를 보장하지는 않지만 비순열(non-permutation) 스케줄에 비해 단순하고 운영이 용이하다는 장점 때문에 실제 생산 환경에서 널리 사용된다. 다수의 기존 연구에서도 순열 스케줄을 가정하여 우수한 결과를 도출하였다(Joo & Kim, 2009; An et al., 2016; Kim & Lee 2019; Lee, 2020; Lee et al., 2021).

작업 $i \in \{1, 2, \dots, n\}$ 는 설비 M1, M2, M3에서 각각 p_{i1} , p_{i2} , p_{i3} 의 공정시간을 갖는다. 앞서 1장에서 소개한 바와 같이 중첩 대기시간 제약에 의해 첫 번째 설비에서 공정이 완료된 후 두 번째 설비에 투입되기까지 대기시간이 a_i 이하여야 하며, 첫 번째 설비에서 공정이 끝난 후 세 번째 설비에 투입되어 공정이 시작되기 전까지 대기시간이 b_i 이내여야 한다. 이를 수식으로 표현하면 아래와 같다.

$$s_{i2} - c_{i1} \leq a_i \quad (1)$$

$$s_{i3} - c_{i1} - p_{i2} \leq b_i \quad (2)$$

식 (1)과 (2)에서 s_{ik} 와 c_{ik} 는 각각 설비 k 에서 작업 i 의 공정시작 시각과 공정종료 시각을 나타낸다. 대기시간에 관한 제약이므로 (2)에서 M2에서 공정이 이루어진 시간인 p_{i2} 는 고려하지 않는다.

특정 작업 스케줄 σ 에서 i 번째로 설비에서 공정이 진행되는 작업을 $[i]$ 로 표현하면 해당 작업의 각 설비에서의 공정종료 시각을 (1)과 (2)에 의해 아래와 같이 표현할 수 있다.

$$c_{[1]1} = p_{[1]1} \quad (3)$$

$$c_{[1]2} = c_{[1]1} + p_{[1]2} \quad (4)$$

$$c_{[1]3} = c_{[1]2} + p_{[1]3} \quad (5)$$

$$c_{[i]1} = \max \left\{ \begin{array}{l} c_{[i-1]1} + p_{[i]1}, \\ c_{[i-1]2} - a_{[i]}, \\ c_{[i-1]3} - b_{[i]} - p_{[i]2} \end{array} \right\}, i = 2, \dots, n \quad (6)$$

$$c_{[i]2} = \max \{ c_{[i-1]2}, c_{[i]1} \} + p_{[i]2}, i = 2, \dots, n \quad (7)$$

$$c_{[i]3} = \max\{c_{[i-1]3}, c_{[i]2}\} + p_{[i]3}, \quad i = 2, \dots, n \quad (8)$$

위의 수식은 중첩 대기시간 제약을 다루었던 기존 연구에서도 제시되었음을 밝힌다 (Kim & Lee, 2019; Lee et al., 2021). 위 수식을 통해 임의의 스케줄이 주어졌을 때, 본 연구에서 다루는 목적식인 총완료시간을 $\sum_{i=1}^n c_{[i]3}$ 와 같이 도출할 수 있다.

Kim & Lee(2019)에서 Makespn 최소화를 위한 MILP 모형이 제안되었으며, 이를 확장하여 총완료시간을 최소화할 수 있는 MILP를 아래와 같이 개발하였다.

$$\text{Minimize } \sum_{i=1}^n c_{[i]3} \quad (9)$$

Subject to

$$\sum_{i=1}^n x_{ij} = 1, j = 1, 2, \dots, n \quad (10)$$

$$\sum_{j=1}^n x_{ij} = 1, i = 1, 2, \dots, n \quad (11)$$

$$s_{[j]k} + \sum_{i=1}^n p_{ik} x_{ij} \leq s_{[j+1]k}, j = 1, 2, \dots, n-1, k = 1, 2, 3 \quad (12)$$

$$s_{[j]k} + \sum_{i=1}^n p_{ik} x_{ij} \leq s_{[j]k+1}, j = 1, 2, \dots, n, k = 1, 2 \quad (13)$$

$$s_{[j]1} + \sum_{i=1}^n (p_{i1} + a_i) x_{ij} \geq s_{[j]2}, j = 1, 2, \dots, n \quad (14)$$

$$s_{[j]1} + \sum_{i=1}^n (p_{i1} + p_{i2} + b_i) x_{ij} \geq s_{[j]3}, j = 1, 2, \dots, n \quad (15)$$

$$c_{[j]3} \geq s_{[j]3} + \sum_{i=1}^n p_{i3} x_{ij}, j = 1, 2, \dots, n \quad (16)$$

x_{ij} 는 결정변수로 작업 i 가 특정 스케줄에서 j 번째 순서로 공정이 진행될 경우 1, 그 외의 경우 0의 값을 갖는다. $s_{[j]k}$ 와 $c_{[j]k}$ 또한 결정변수로서 j 번째 순서로 공정이 진행되는 작업의 k 번째 설비에서의 공정시작 시각과 공정종료 시각을 각각 나타낸다. (9)는 목적식이 총완료시간 최소화임을 나타내며, (10)과 (11)은 특정 작업이 특정 순서에 한 번씩만 진행되도록 강제한다. (12)는 특정 설비에서 연속으로 진행되는 작업들의 공정시작 시각 관계를 나타내며, (13)은 특정 작업에 대해 연속된 설비 사이에서 공정시작 시각의 관계를 나타낸다. (14)와 (15)는 중첩 대기시간 제약을 만족하도록 하는 제약식이다. 마지막으로 (16)은 마지막 설비에서 각 작업의 종료시각을 계산하는 데 활용되며 이를 통해 총완료시간을 도출하여 최소화할 수 있다.

제시된 MILP는 본 연구에서 다루고자 하는 문제의 특성을 수학적으로 명확히 모델링하며, 작은 크기의 문제들에 대해 합리적인 계산시간 내에 최적해를 도출한다. 이에 대한 자세한 실험 결과는 5장에서 제시한다.

4. 휴리스틱 알고리즘

본 장에서는 유전알고리즘과 지역탐색 기법을 활용한 휴리스틱을 제안한다. 유전알고리즘은 스케줄링 문제에 널리 적용되어 왔으며, 우수한 성능을 보임을 기존 연구들을 통해 알 수 있다(Eroglu et al., 2014; Tseng & Lin, 2009; Vallada & Ruiz, 2011; Lee, 2020; Lee et al., 2021). 유전알고리즘의 세부 구성요소들은 실제 생물의 진화 과정에서 영감을 얻어 개발되었으며 구체적으로 선택, 교차, 변이 등의 과정을 포함한다. 본 연구에서 제안하는 휴리스틱의 세부 사항들은 다음의 세부 절들을 통해 설명한다.

4.1 유전자 표현방법 및 적합도

앞서 3장에서 소개한 바와 같이 본 연구에서는 순열 플로우샵을 가정한다. 즉, 첫 번째 설비 M1에서의 작업 투입 순서만 결정되면 M2와 M3에서도 동일한 순서로 공정이 이루어진다. 또한 작업순서만 결정이 되면 수식 (3)-(8)을 통해 대기시간 제약을 만족하는 스케줄 하에서 각 작업의 공정종료 시각을 도출할 수 있다. 따라서 결정되어야 하는 사항은 M1에서의 작업 순서이다. 이를 바탕으로 유전자가 표현해야 하는 정보는 작업 순서임을 알 수 있으며, 따라서 아래와 같은 유전자 표현방법을 사용한다.

1	4	2	5	3
---	---	---	---	---

Fig. 1 Gene Representation

<Fig. 1>은 유전자 표현의 예를 나타내며, 각 설비에 1, 4, 2, 5, 3의 작업 순서로 공정이 이루어져야 함을 의미한다. 본 연구의 목적식은 총완료시간 최소화이므로 유전자의 적합도는 유전자에 표현된 작업순서로 작업을 했을 때 얻게 되는 총완료시간으로 정의할 수 있다.

4.2 초기모집단(initial population)

유전알고리즘의 중요한 특징 중 하나는 초기모집단

을 구성하는 유전자들을 바탕으로 교차, 변이 등을 통해 진화된 유전자들을 얻어간다는 점이다. 따라서 일반적으로 초기모집단을 구성하는 유전자들이 우수할 수록 더 빠른 시간 안에 최적해에 근접한 해를 찾을 확률이 높아진다고 할 수 있다. 본 연구에서는 우수한 초기모집단을 구성하기 위해 기존 연구들을 참고하여 아래 <Table 1>의 8개 알고리즘을 사용하였다. 구체적으로 아래 소개된 8개의 알고리즘을 이용하여 초기모집단을 구성하고 부족한 유전자들은 무작위(random) 방식으로 생성하였다.

Table 1. Algorithms for Initial Population

Algorithm	Description
A1	SPT with p_{i1}
A2	SPT with p_{i2}
A3	SPT with p_{i3}
A4	SPT with $\sum_{k=1}^3 p_{ik}$
A5	Algorithm by Rajendran & Chaudhuri(1991)
A6	Algorithm by Aldowaisan & Allahverdi(2004)
A7	Algorithm by Laha & Sarin(2009)
A8	Algorithm by Gao et al.(2013)

단일설비 스케줄링에서는 SPT(Shortest Processing Time) 규칙을 이용하여 총완료시간을 최소화할 수 있음이 알려져 있다. 이에 착안하여 <Table 1>의 A1-A4가 적용되었다. A1-A4는 각각 작업별 첫 설비에서의 공정시간, 두 번째 설비에서의 공정시간, 세 번째 설비에서의 공정시간, 첫 번째부터 세 번째 공정설비들에서의 공정시간의 총합을 기준으로 SPT로 작업순서를 결정하며 결정된 순서는 <Fig. 1>과 같이 유전자로 표현된다. A5-A8은 문헌 연구를 통해 발견하였으며, 모두 플로우샵에서 총완료시간 최소화를 위해 개발된 휴리스틱 알고리즘들이다. 단, 세부적인 가정사항들은 본 연구에서 다루는 문제와 다소 차이가 있다. 예를 들면 Aldowaisan & Allahverdi(2004)의 경우 공정간 대기시간을 허용하지 않는 no-wait 플로우샵을 가정한다. 앞서 소개한 바와 같이 본 연구에서 다루는 문제는 a_i 와 b_i 에 따라 no-wait 플로우샵 문제가 될 수도 있고 대기시간 제약이 없는 문제가 될 수도 있는 포괄적인 문제이다. 따라서 특정 환경에서는 Aldowaisan & Allahverdi(2004)의 알고리즘으로 도출한 해가 우수한 성능을 보일 가능성이 있기에 다양한 알고리즘들을 초기해 구성에 포함하였다.

4.3 선택(Selection)

유전알고리즘은 현재 세대의 유전자들을 바탕으로 교차, 변이 등을 통해 새로운 유전자들을 생성하고 추가한다. 이 후 ‘선택’ 과정을 통해 현재 세대의 유전자들 중 일부가 다음 세대로 전달된다. 선택과정에서 다양한 방법들을 활용할 수 있다(Moon, 2003; Lee, 2020; Lee et al., 2021). Lee(2020)는 대기시간 제약이 있는 2-설비 플로우샵 스케줄링 문제에서 룰렛 휠 방식이 우수한 성능을 보임을 실험적으로 증명하였다. 이러한 기존 연구결과를 바탕으로 본 논문에서는 룰렛 휠 방식을 사용한다. 룰렛 휠 방식은 기본적으로 우수한 유전자 즉, 본 연구에서 다루는 문제의 경우 작은 총완료시간을 갖는 유전자에 다음 세대에 선택될 높은 확률을 부여한다. 반면 비교적 큰 총완료시간을 갖는 유전자에도 일정 확률을 부여하여 다양성을 유지한다. 본 연구에서는 각각의 유전자에 확률을 부여하는 방식을 아래와 같이 정의하였다.

$$P_i = \frac{f^{\max} - f_i}{\sum_j (f^{\max} - f_j)} \quad (17)$$

P_i 는 유전자 i 가 다음 세대로 전달될 확률이며, f_i 는 앞서 4장 1절에서 소개한 적합도 즉, 총완료시간을 의미한다. $f^{\max}$ 는 현 세대의 모든 유전자들이 갖는 적합도 중 최대값을 의미한다. 본 연구에서 고려하는 목적식은 총완료시간 최소화이므로 적합도(즉, 총완료시간)가 낮을수록 높은 확률을 부여해야 한다. 이를 위해 f_i 가 아닌 $f^{\max} - f_i$ 를 사용한다. 즉 현재 세대의 유전자 적합도 중 가장 큰 값과 특정 유전자의 적합도 차이가 클수록 높은 확률을 부여 받는다.

4.4 교차(Crossover)

‘교차’는 두 부모 유전자의 조합을 통해 새로운 자식 유전자를 만들어 내는 과정이다. 교차를 위해 일점(one-point) 교차, 이점(two-point) 교차 등 다양한 방법들이 사용된다(Moon, 2003; Lee, 2020). Hosseinabadi et al. (2019)과 Hasancebi & Erbatur(2000)은 유전알고리즘을 위해 사용되는 다양한 교차 방법을 적용하고 비교실험을 통해 스케줄링 문제에서 일점 교차가 우수한 성능을 보임을 확인하였다. Lee(2020)에서도 일점 교차 방식이 대기시간 제약이 있는 2-설비 플로우샵 스케줄링 문제에서 비교적 우수한 성능을 보임을 확인할 수 있다. 이러한

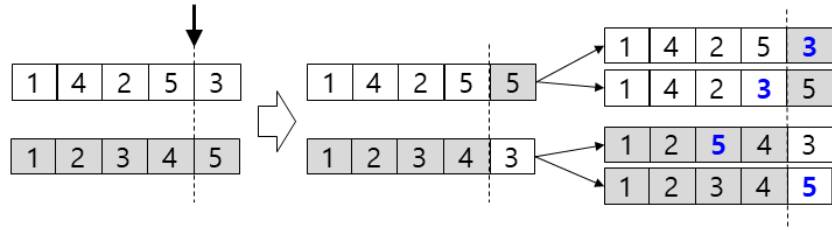


Fig. 2 One-point Crossover

기존 연구 결과를 바탕으로 본 연구에서는 일점 교차 방식을 적용한다.

〈Fig. 2〉는 본 연구에서 사용하는 일점 교차 방식을 나타낸다. 왼쪽에 표현된 바와 같이 일점 교차에서는 우선 교차가 일어날 한 지점을 임의로 선택하게 된다. 선택된 지점을 기준으로 두 부모 유전자의 교차가 일어난다. 〈Fig. 2〉에서 볼 수 있는 바와 같이 교차 후 두 개의 자식 유전자가 생성됨을 알 수 있으며, 각각 (1, 4, 2, 5, 5)와 (1, 2, 3, 4, 3)을 작업순서로 갖는다. 본 예제를 통해 쉽게 알 수 있듯이 일점 교차의 경우 교차가 일어난 후 특정 작업들이 겹치거나 혹은 특정 작업이 누락될 수 있다. (1, 4, 2, 5, 5)에서는 작업 5가 두 번 등장하며, (1, 2, 3, 4, 3)에서는 작업 5가 누락되었다. 간단한 보정 작업을 통해 이러한 문제를 해결할 수 있다. 〈Fig. 2〉의 오른쪽에 표현된 바와 같이 선택된 점의 왼쪽 작업들은 고정하고 오른쪽 작업들을 보정해주거나, 유사하게 오른쪽 작업들은 고정하고 왼쪽 작업들을 보정해줄 수 있다. 〈Fig. 2〉의 붉은 파란색으로 표현된 작업들은 보정된 작업을 나타낸다. 전술한 보정 방식을 적용하면 2개의 부모 유전자로부터 〈Fig. 2〉와 같이 총 4개의 자식 유전자가 생성된다. 생성된 유전자들은 모두 모집단에 포함시킨다. 이를 일반적인 알고리즘으로 표현하면 다음과 같다.

Step 0: p 는 1과 $n-1$ 사이에서 무작위로 생성된 정수이며, π_1 과 π_2 는 부모유전자 1과 2가 나타내는 작업순서, $\pi_i(j, k)$, $i = 1, 2$, $1 \leq j \leq k \leq n$, 는 π_i 의 j 번째부터 k 번째 작업을 포함하는 부분 작업순서를 나타낸다.

Step 1: $\pi_1(1, p)$ 와 $\pi_2(p+1, n)$ 그리고 $\pi_2(1, p)$ 와 $\pi_1(p+1, n)$ 가 결합된 작업순서를 각각 π_a^c 와 π_b^c 라 하자.

Step 2: π_a^c 의 1부터 p 번째 작업 중 $p+1$ 번째부터 n 번째 작업과 겹치는 작업이 있을 경우 해당 작업을 임의의 다른 작업으로 변경하여 겹치지 않도록 한다. 변경 완료된 작업순서를 π_1^c 라 한다.

Step 3: π_a^c 의 $p+1$ 번째부터 n 번째 작업 중 1부터 p 번째 작업과 겹치는 작업이 있을 경우 해당 작업을 임의의 다른 작업으로 변경하여 겹치지 않도록 한다. 변경 완료된 작업순서를 π_2^c 라 한다.

Step 4: π_b^c 의 1부터 p 번째 작업 중 $p+1$ 번째부터 n 번째 작업과 겹치는 작업이 있을 경우 해당 작업을 임의의 다른 작업으로 변경하여 겹치지 않도록 한다. 변경 완료된 작업순서를 π_3^c 라 한다.

Step 5: π_b^c 의 $p+1$ 번째부터 n 번째 작업 중 1부터 p 번째 작업과 겹치는 작업이 있을 경우 해당 작업을 임의의 다른 작업으로 변경하여 겹치지 않도록 한다. 변경 완료된 작업순서를 π_4^c 라 한다.

Step 6: π_1^c , π_2^c , π_3^c , π_4^c 가 새롭게 생성된 자식유전자의 작업순서들을 나타내며 이들을 모집단에 포함시킨다.

마지막으로 교차와 관련된 파라미터로는 교차 비율(0에서 1사이의 값)이 있다. 교차비율에 따라 (모집단의 유전자 수 $\times$ 교차 비율) 만큼의 교차가 일어나며, 매 교차 시에 두 개의 부모 유전자를 모집단에서 무작위로 선정하는 방식을 사용한다.

4.5 변이(Mutation)

모집단에 있는 모든 유전자를 대상으로 난수(0에서 1 사이)를 생성하고 해당 값이 주어진 변이 확률보다 작을 경우 변이를 진행한다. 기존 유전자와 변이된 유전자 모두 모집단에 포함시킨다. 본 연구에서는 변이 방식으로 삽입(insertion)과 교환(interchange) 두 가지 방식을 사용하며, 두 방식 모두 변이를 위해 널리 사용되고 우수한 성능을 보임이 알려져 있다(Lee, 2020). 삽입의 경우 무작위로 하나의 작업과 해당 작업이 삽입될 위치를 선정하여 새로운 위치에 삽입하는 방식이다. 교환의 경우 유전자가 가지는 작업 순서에서 두 작업을 무작위로 선정하여 해당 작업들의 위치를 교환하는 방식이다.

4.6 지역탐색(Local Search)

휴리스틱 알고리즘의 성능 향상을 위해 추가적으로 지역탐색 기법을 활용한다. 지역탐색으로 인해 계산 시간이 급격히 증가하는 것을 방지하기 위해 매 세대별 적합도 즉, 총완료시간을 기준으로 우수한 10% 유전자에 대해서만 지역탐색이 적용되도록 디자인하였다. 일반적으로 지역탐색을 위해 많이 사용되는 방법은 앞서 4장 5절에서도 소개하였던 삽입과 교환방법이 있다(Tseng & Lin, 2009). 본 연구에서도 지역탐색을 위해 삽입과 교환 방식을 사용했으며, 구체적인 적용 방식은 아래와 같다.

4.6.1 삽입

현재 유전자가 담고 있는 작업순서 정보를 바탕으로 모든 작업들에 대해 현재 위치가 아닌 다른 위치에 삽입해 보고 가장 작은 총완료시간을 주는 새로운 작업순서를 도출하여 해당 유전자의 정보를 변경한다.

Step 0: $TCT^* = \infty$, $i = 1$, $j = 2$, π = 유전자가 나타내는 작업순서

Step 1: $i = j$ 이면, Step 2로 이동한다. 그렇지 않으면, π 의 i 번째 작업이 j 번째 작업이 되도록 옮겨 삽입한다. 이를 통해 얻은 작업순서를 π' 라 하자. π' 로부터 도출된 총완료시간이 TCT^* 보다 작으면 TCT^* 를 해당 값으로 업데이트 하고 π^* 를 π' 로 업데이트 한다.

Step 2: $j = n$ 이면 Step 3으로 이동한다. 그렇지 않으면 $j = j + 1$ 로 업데이트하고 Step 1로 이동한다.

Step 3: $i = n$ 이면 Step 4로 이동한다. 그렇지 않으면 $i = i + 1$, $j = 1$ 로 업데이트하고 Step 1로 이동한다.

Step 4: 유전자의 작업순서를 π^* 로 업데이트하고 종료한다.

4.6.2 교환

현재 유전자가 담고 있는 작업순서 정보를 바탕으로 모든 작업쌍(pair)에 대해 위치를 교환해보고 가장 작은 총완료시간을 주는 새로운 작업순서를 도출 후 해당 작업순서로 유전자의 정보를 변경한다. 구체적인 적용방식은 아래와 같다.

Step 0: $TCT^* = \infty$, $i = 1$, $j = 2$, π = 유전자가 나타내는 작업순서

Step 1: π 의 i 번째 작업과 j 번째 작업의 위치를 서로 바꾼다. 이를 통해 얻은 작업순서를 π' 라 하자. π' 로부터 도출된 총완료시간이 TCT^* 보다 작으면 TCT^* 를 해당 값으로 업데이트 하고 π^* 를 π' 로 업데이트 한다.

Step 2: $j = n$ 이면 Step 3으로 이동한다. 그렇지 않으면 $j = j + 1$ 로 업데이트하고 Step 1로 이동한다.

Step 3: $i = n - 1$ 이면 Step 4로 이동한다. 그렇지 않으면 $i = i + 1$ 로 업데이트 하고 업데이트 된 i 로 $j = i + 1$ 로 업데이트 한 후 Step 1로 이동한다.

Step 4: 유전자의 작업순서를 π^* 로 업데이트하고 종료한다.

4.7 제안된 방법의 흐름도

제안된 휴리스틱 방법은 미리 지정된 최대 세대수 G 에 도달하면 종료되며, 전체적인 흐름도는 <Fig. 3>과 같다.

흐름도를 간단히 요약하면 우선 본장 2절에서 소개된 8개의 알고리즘들과 무작위 방법을 통하여 초기모집단을 구성한다. 초기모집단이 구성되고 나면 앞서 소개한 교차와 변이를 통해 모집단에 새로운 유전자를 추가하고, 지역탐색 기법을 통하여 해의 성능을 향상시킨다. 이후 선택과정을 통해 선별된 유전자들을 다음세대에 전달한다. 세대 진화 횟수 g 가 주어진 G 에 도달할 때까지 과정을 반복한다.

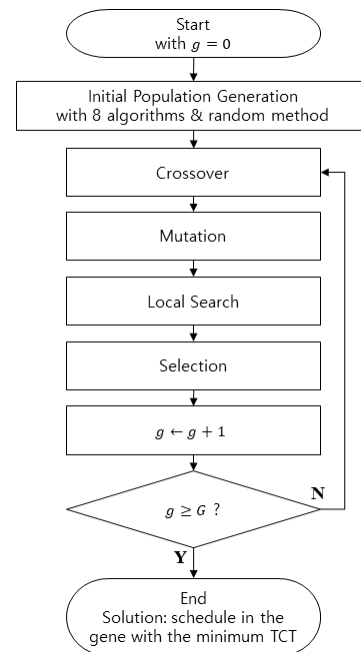


Fig. 3 Flowchart of the Proposed Algorithm

5. 실험결과

제안된 휴리스틱 알고리즘은 Java로 구현하였고 모든 실험은 PC(Intel Core i7-9700, 32GB RAM)에서 진행하였다.

5.1 파라미터 설정

전술한 바와 같이 본 연구에서 제안한 휴리스틱 알고리즘은 유전알고리즘을 기반으로 한다. 유전알고리즘의 대표적인 파라미터로는 모집단의 크기(N), 세대수(G), 교차 비율(c), 변이 확률(m)이 있다. 일반적으로 고려하는 문제 특성에 적합한 파라미터들을 선정하기 위해 예비실험을 활용한다. 본 연구에서는 기존 연구들을 참고하여 다음과 같은 파라미터 후보들을 고려하였다: $N \in \{n, 2n, \dots, 10n\}$ (n 은 작업의 수), $G \in \{100, 200, \dots, 1000\}$, $c \in \{0.3, 0.5, 0.7, 0.9\}$, $m \in \{0.05, 0.1, 0.2\}$ (Eroglu et al., 2014; Tseng & Lin, 2009; Lee et al., 2021). 전술한 파라미터 후보들에 대해 예비실험을 진행했으며 이 때 작업의 수 n 은 50으로 고정하였고 공정시간은 1에서 50 사이의 정수를 무작위로 생성하였다. 중첩 대기시간 제약을 나타내는 a_i 와 b_i 는 공정시간을 고려하여 다음과 같이 무작위로 생성하였다: $a_i \in \{0, 1, \dots, 50\}$, $b_i \in \{a_i, a_i + 1, \dots, 100\}$.

동시에 모든 조합을 고려할 경우 총 $1200 (= 10 \times 10 \times 4 \times 3)$ 개의 경우의 수가 있어, Lee et al.(2021)을 참고하여 다른 파라미터 값들은 기존 연구에서 사용된 값으로 고정하고 $N \rightarrow G \rightarrow c \rightarrow m$ 순서로 순차적으로 예비실험을 진행하여 최소의 총완료시간을 주는 파라미터 값들을 선정하였다. 최종 선정된 파라미터는 다음과 같다: $N = 3n$, $G = 300$, $c = 0.5$, $m = 0.2$. 자세한 예비실험 결과는 본 연구의 주된 내용이 아니므로 소개를 생략하고, 대안에 따른 총완료시간 차이가 미미한 경우에는 알고리즘의 복잡도를 낮추는 방향으로 파라미터를 선정하였음을 밝힌다.

5.2 실험세팅

본 연구에서 고려되는 주된 요소들은 작업의 수, 공정시간, 대기시간 제약이다. 작업 수에 따른 알고리즘 성능과 MILP로부터 도출된 해와의 비교를 위해 작업의 수가 10개, 20개, 30개, 40개, 50개인 경우를 고려하였다. 공정시간은 1에서 50 사이의 정수를 무작위로 생성하였다. 앞서 소개한 바와 같이 본 연구에서 다루는 문제는 대기시간 제약 측면에서 제약이 엄격할 경우 no-wait 플로우샵과 유사해지고, 반대로 제약이 느슨할 경우 대기시간에 제약이 없는 일반적인 플로우샵과 유사하다. 따라서 이러한 점을 모두 고려하기 위해 공정시간 범위를 바탕으로 <Table 2>와 같이 세 가지 유형의 대기시간 제약을 고려한다. 중첩 대기시간 제약을 다룬 기존 연구에서도 유사한 실험 세팅이 사용되었다(Kim & Lee, 2019; Lee et al., 2021).

5.3 초기모집단 구성을 위한 알고리즘 및 지역탐색 기법 효과성 검증

<Table 3>은 초기모집단 구성을 위한 알고리즘 및 지역탐색 기법의 효과성을 검증하기 위한 실험 결과를 나타낸다. <Table 3>의 HA1, HA2, HA3의 의미는 다음과 같다.

- HA3: 제안하는 최종 알고리즘 (4장에서 소개된 모든 방법 포함)
- HA2: HA3에서 지역탐색 기법(4장 6절 참고) 제외
- HA1: HA2에서 초기모집단 구성을 위한 알고리즘(4장 2절 참고) 제외 즉, 초기모집단을 무작위 방법으로 구성

<Table 3>의 첫 번째 열은 앞서 본장 2절에서 제안된 실험세팅을 의미하며 두 번째 열은 작업의 수를 나타낸다. 각각의 조합에 대해 20개의 인스턴스를 무작위로 생성하였고 각각에 대해 HA1-3을 적용하여 평균 총완료시간

Table 2. Experimental Settings

Processing Time Range	Waiting Time Limits	Experimental Setting
$p_{ij} \in \{1, 2, \dots, 100\}$	$a_i \in \{0, 1, \dots, 10\}$, $b_i \in \{a_i, a_i + 1, \dots, 20\}$	(1)
	$a_i \in \{0, 1, \dots, 50\}$, $b_i \in \{a_i, a_i + 1, \dots, 100\}$	(2)
	$a_i \in \{0, 1, \dots, 100\}$, $b_i \in \{a_i, a_i + 1, \dots, 200\}$	(3)

간(TCT)과 초 단위의 계산시간(CT)을 <Table 3>에 기입하였다. 유전알고리즘의 특성상 같은 인스턴스 하에서도 매번 도출되는 해가 달라질 수 있으므로 신뢰도 있는 결과를 위해 HA1-3을 각 인스턴스에 대해 10회씩 반복 적용하여 평균값을 활용하였다. 마지막 두 열은 각각 HA2와 HA1 그리고 HA3와 HA2로부터 도출된 해의 목적식 차이를 %로 나타낸다. 전반적으로 모든 경우에서 HA1에서 HA3으로 갈수록 총완료시간(목적식)의 값이 감소함을 알 수 있다. 대기시간 제약에 관한 실험세팅 측면에서 보면 대기시간 제약이 엄격할수록 즉, 세팅(3)에서 (1)로 갈수록 HA2와 HA3에 의한 해의 개선이 증가함을 알 수 있다. 일반적으로 시간제약이 엄격할수록 문제의 난이도가 높아진다는 점에서, 어려운 문제일수록 HA2와 HA3에 의한 해의 개선이 증가한다고 해석할 수 있다.

HA1과 HA2를 비교해보면 작업의 수 n 이 커질수록 도출된 총완료시간의 차이가 커짐을 알 수 있다. 이는 작업의 수가 작은 경우 단순한 유전알고리즘 만으로도 우수한 해를 얻을 수 있으나 작업의 수가 커질수록 문제가 어려워지고 본 논문에서 제안한 추가적인 방법들을 통해 알고리즘의 성능 개선이 필요하다고 해석할 수 있다. 예를 들어 실험세팅(1) 하에서 $n=10$ 인 경우 HA2가 HA1에 비해 0.6% 개선된 총완료시간을 주지만, $n=50$ 인 경우 5.2%의 개선이 일어남을 알 수 있다. 실험세팅 (2)와 (3)에서도 유사한 패턴을 관측할 수 있다. 이를 통해 본 연구에서 초기모집단 구성을 위해 사용한 알고리즘들이 해의

개선에 긍정적인 효과를 나타냄을 알 수 있다. HA3와 HA2의 비교를 통하여 제안된 지역탐색 기법의 효과성을 확인할 수 있다. 앞선 HA2와 HA1의 비교와 유사하게 작업의 수가 커질수록 도출된 총완료시간의 차이가 커지는 경향을 보인다. 예를 들어 실험세팅 (1) 하에서 $n=10$ 인 경우 HA3가 HA2에 비해 0.3% 개선된 총완료시간을 주지만, $n=50$ 인 경우 2.1%의 개선이 일어난다. 실험세팅(2)와 (3)에서도 유사한 경향성을 확인할 수 있다.

계산시간 측면에서 살펴보면 HA1에서 HA3으로 갈수록 계산시간이 증가함을 알 수 있다. 하지만 작업이 50개 일 때도 HA3를 통해 평균 15초 이내로 해를 도출할 수 있음을 알 수 있다. 계산시간과 관련해서는 추가적으로 큰 n 을 고려하여 향후 절에서 자세히 논의한다.

5.4 MILP와의 비교를 통한 성능 검증

<Table 4>는 본 논문에서 최종적으로 제안하는 알고리즘인 HA3와 본 논문의 3장에서 소개된 MILP의 비교결과를 보여준다. 각각의 방법에 대해 20개 인스턴스의 평균 총완료시간과 계산시간을 <Table 4>에 기입하였다. MILP를 풀기위해 Gurobi (Version 9.1.2)를 사용했으며 1시간의 계산시간 제약을 두고 실험을 진행하였다. 1시간 이내에 최적해가 도출되면 해당 최적해의 목적식 값을 이용했고, 그렇지 않은 경우 1시간 경과된 시점에서 현재의 최소 목적식(current best solution) 값을 사용했다.

Table 3. Comparison between HA1, HA2, and HA3

Exp. Setting	n	HA1		HA2		HA3		Gap(%) HA2 vs. HA1	Gap(%) HA3 vs. HA2
		TCT	CT	TCT	CT	TCT	CT		
(1)	10	3822.7	0.02	3798.1	0.02	3786.4	0.16	0.6	0.3
	20	11909.2	0.03	11714.6	0.03	11554.5	1.10	1.6	1.4
	30	26328.2	0.03	25586.6	0.05	25222.0	3.01	2.8	1.4
	40	45060.1	0.04	43157.5	0.08	42483.1	6.91	4.2	1.6
	50	70123.7	0.05	66485.6	0.13	65111.9	13.38	5.2	2.1
(2)	10	3502.0	0.02	3497.3	0.02	3481.8	0.14	0.1	0.4
	20	11271.0	0.03	11162.6	0.03	11034.0	1.00	1.0	1.2
	30	24486.8	0.04	23913.2	0.05	23580.2	3.13	2.3	1.4
	40	41256.6	0.05	39933.4	0.08	39290.2	7.14	3.2	1.6
	50	62470.2	0.06	59735.8	0.13	58794.1	13.70	4.4	1.6
(3)	10	3567.2	0.02	3560.9	0.02	3547.9	0.15	0.2	0.4
	20	11068.6	0.03	10980.1	0.03	10846.4	0.99	0.8	1.2
	30	23561.9	0.04	23154.1	0.05	22870.1	3.14	1.7	1.2
	40	41884.7	0.05	40801.2	0.08	40259.4	7.17	2.6	1.3
	50	62691.3	0.06	60618.6	0.14	59763.4	14.20	3.3	1.4

Table 4. Comparison between HA3 and MILP

Exp. Setting	n	HA3		MILP			Gap(%) HA3 vs. MILP
		TCT	CT	TCT	CT	INO(%)	
(1)	10	3786.4	0.16	3785.5	0.13	0	0.02
	20	11554.5	1.10	11506.7	82.98	0	0.42
	30	25222.0	3.01	25045.4	3461.22	95	0.71
	40	42483.1	6.91	42360.2	3600.00	100	0.29
	50	65111.9	13.38	65154.6	3600.00	100	-0.06
(2)	10	3481.8	0.14	3478.3	0.11	0	0.10
	20	11034.0	1.00	10996.1	23.49	0	0.36
	30	23580.2	3.13	23349.1	1854.61	35	0.98
	40	39290.2	7.14	38952.4	3435.79	95	0.85
	50	58794.1	13.70	58444.9	3600.00	100	0.61
(3)	10	3547.9	0.15	3547.5	0.11	0	0.01
	20	10846.4	0.99	10807.1	13.27	0	0.36
	30	22870.1	3.14	22673.3	2272.00	50	0.87
	40	40259.4	7.17	39960.5	3444.52	95	0.75
	50	59763.4	14.20	59326.2	3600.00	100	0.75

<Table 4>의 7열 INO(%)는 1시간 내에 MILP를 통해 최적해가 도출되지 않은 인스턴스의 비율을 의미한다. 마지막 열은 HA3와 MILP로부터 도출된 목적식 값의 차이를 %로 나타낸다.

우선 MILP의 계산시간을 살펴보면 실험세팅(1)~(3)에 대해 작업의 수가 40개 이상이 되면 거의 모든 경우 1시간 내에 최적해를 도출하지 못함을 알 수 있다. 특히 실험세팅(1) 즉, 시간제약이 엄격한 경우 MILP의 계산시간이 길어짐을 알 수 있다. 반면 HA3의 경우 계산시간이 실험세팅에 영향을 받지 않고, 작업이 50개인 경우도 15초 이내에 해를 도출함을 알 수 있다. 목적식 값의 차이를 살펴보면 모든 경우 HA3로부터 도출된 목적식 값이 짧은 계산시간에도 불구하고 MILP로부터 도출된 목적식 값과 1% 이내의 차이가 남을 알 수 있다. 특히 실험세팅(1), $n=50$ 인 경우 HA3의 계산시간은 MILP의 0.37%에 불과하지만 오히려 더 우수한 해를 도출함을 알 수 있다.

5.5 n 이 큰 경우 HA3의 계산시간

본 절에서는 n 이 큰 경우 HA3의 계산시간에 대한 실험 결과를 소개한다. 구체적으로 $n \in \{100, 200, 300\}$ 인 경우를 고려하였으며, 각각 실험세팅(3) 하에서 10개의 인스턴스를 무작위로 생성하여 평균 계산시간(초)을 도출하고 <Table 5>에 기입하였다. n 이 커질수록 지역탐색 등으로 인해 계산시간이 급격히 증가하지만, $n=300$ 인

경우에도 1시간 이내에 해를 도출할 수 있다. 동일한 인스턴스에 대해 12시간의 시간 제약을 두고 MILP를 적용해 보았으나 모든 경우 12시간 이내에 최적해를 도출할 수 없었음을 밝힌다.

Table 5. Computation Time of HA3 with Large n

n	HA3 Computation Time (s)
100	104.48
200	913.14
300	3181.71

6. 결 론

본 논문에서는 총완료시간 최소화를 위해 중첩 대기 시간 제약을 갖는 3-설비 플로우샵 스케줄링 문제를 다루었다. 최적해 도출을 위한 MILP를 소개하였고, 빠른 시간 안에 최적해에 근사한 해를 얻을 수 있는 휴리스틱 방법 또한 제안하였다. 실험을 통하여 본 연구에서 제안한 방법들에 대해 각각의 효과성을 입증하였다. 또한 제안된 휴리스틱 알고리즘이 빠른 시간 내에 최적해와 1% 미만의 차이를 가지는 우수한 해를 도출할 수 있음을 실험을 통해 확인하였다.

향후 연구는 크게 세 가지 방향으로 생각해 볼 수 있다. 첫째로, 총완료시간 외에 다양한 목적식을 고려하는 방

향이다. 본 연구에서는 각 작업의 납기를 고려하지 않았으나 납기를 고려하여 지연시간 최소화, 지연 작업 수 최소화 등을 목적식으로 고려할 수 있다. 둘째로, 더욱 일반적인 생산시스템에 대한 문제를 생각할 수 있다. 본 연구에서는 3-설비 플로우샵 스케줄링 문제를 다루었으나 일반적인 m-설비 그리고 더욱 다양한 중첩 대기시간 제약 환경을 고려하여 연구결과를 확장할 수 있을 것으로 생각된다. 마지막으로, 알고리즘의 계산 효율성을 높이는 방향으로 연구를 진행 할 수 있다. 비록 제안된 휴리스틱 알고리즘이 현실적인 문제를 비교적 빠른 시간 안에 해결 가능 하지만, 실험을 통해 보인 바와 같이 작업의 수가 커질수록 계산시간이 급격히 증가한다. 계산 부하를 줄이면서도 우수한 해를 얻을 수 있는 방법을 개발하는 것 또한 향후 좋은 연구 방향이 될 수 있다고 판단된다.

REFERENCES

- [1] Aldowaisan, T. & Allahverdi, A.(2004). New heuristics for m-machine no-wait flowshop to minimize total completion time, *Omega*, 32(5), 345-352.
- [2] An, Y.-J., Kim, Y.-D., & Choi, S.-W.(2016). Minimizing makespan in a two-machine flowshop with a limited waiting time constraint and sequence-dependent setup times. *Computers & Operations Research*, 71, 127-136.
- [3] Eroglu, D. Y., Ozmutlu, H. C., & Ozmutlu, S.(2014). Genetic algorithm with local search for the unrelated parallel machine scheduling problem with sequence-dependent set-up times. *International Journal of Production Research*, 52(19), 5841-5856.
- [4] Framinan, J. M. & Leisten, R.(2003). An efficient constructive heuristic for flowtime minimization in permutation flow shops. *Omega*, 31(4), 311-317.
- [5] Gao, K., Pan, Q., Suganthan, P. N., & Li, J.(2013). Effective heuristics for the no-wait flow shop scheduling problem with total flow time minimization. *International Journal of Advanced Manufacturing Technology*, 66, 1563-1572.
- [6] Hasancebi, O. & Erbatur, F.(2000). Evaluation of crossover techniques in genetic algorithm based optimum structural design. *Computer & Structures*, 78, 435-448.
- [7] Hong, J. H. & Yoon S.-H.(2021). To minimize the weighted number of early and tardy jobs in a two-machine flow shop. *Journal of the Korean Society of Supply Chain Management*, 21(1), 69-78.
- [8] Hosseinabadi, A. A. R., Vahidi, J., Saemi, B., Sangaiah, A. K., & Elhoseny, M.(2019). Extended genetic algorithm for solving open-shop scheduling problem. *Soft Computing*, 23, 5099-5116.
- [9] Joo, B.-J. & Kim, Y.-D.(2009). A branch-and-bound algorithm for a two-machine flowshop scheduling problem with limited waiting time constraints. *Journal of the Operational Research Society*, 60, 572-707.
- [10] Ju, F., Li, J., & Horst, J. A.(2017). Transient analysis of serial production lines with perishable products: Bernoulli reliability model. *IEEE Transactions on Automatic Control*, 62(2), 694-707.
- [11] Kim, H.-J. & Lee, J.-H.(2019). Three-machine flow shop scheduling with overlapping waiting time constraints. *Computers & Operations Research*, 101, 93-102.
- [12] Laha, D. & Sarin, S. C.(2009). A heuristic to minimize total flow time in permutation flow shop. *Omega*, 37(3), 734-739.
- [13] Lee, J.-Y.(2020). A genetic algorithm for a two-machine flowshop with a limited waiting time constraint and sequence-dependent setup times. *Mathematical Problems in Engineering*, 2020, 1-13.
- [14] Lee, J.-H, Yu, T.-S., & Park, K.(2021). Scheduling of flow shop with overlapping waiting time constraints using genetic algorithm. *Journal of the Korean Institute of Industrial Engineers*, 47(1), 34-44.
- [15] Lee, J.-H, Zhao, C., Li, J., & Papadopoulos, C. T.(2018). Analysis, design, and control of Bernoulli production lines with waiting time constraints. *Journal of Manufacturing Systems*, 46, 208-220.
- [16] Moon, B.-R.(2003). *Genetic Algorithm*. Dooyang-Sa, Seoul, Korea.
- [17] Pinedo, M. L.(2016). *Scheduling: Theory, Algorithms, and Systems* (5th Edition), Springer, New York.
- [18] Rajendran, C. & Chaudhuri, D.(1992). An efficient heuristic approach to the scheduling of jobs in a flowshop. *European Journal of Operational Research*, 61(3), 318-325.
- [19] Tseng, L.-Y. & Lin, Y.-T.(2009). A hybrid genetic local search algorithm for the permutation flowshop scheduling problem. *European Journal of Operational Research*, 198, 84-92.
- [20] Vallada, E. & Ruiz, R.(2011). A genetic algorithm for the unrelated parallel machine scheduling problem with

sequence dependent setup times. *European Journal of Operational Research*, 211, 612-622.

- [21] Wang, J., Hu, Y., & Li, J.(2010). Transient analysis to design buffer capacity in dairy filling and packing production lines. *Journal of Food Engineering*, 98(1), 1-12.

- [22] Yu, T.-S., Kim, H.-J., & Lee, T.-E.(2017). Minimization of waiting time variation in a generalized two-machine flowshop with waiting time constraints and skipping jobs. *IEEE Transactions on Semiconductor Manufacturing*, 30(2), 155-165.



이 준 호

KAIST 산업및시스템공학 박사
현재: 충남대학교 경영학부 조교수
관심분야: 스케줄링, SCM



김 현 정

KAIST 산업및시스템공학 박사
현재: KAIST 산업및시스템공학과
조교수
관심분야: 스케줄링, SCM

빅데이터 DEA를 위한 기계학습의 적용*

응위엔투이즈영** · 임성묵**†

**동국대학교-서울캠퍼스 경영대학 경영학과

Application of Machine Learning to DEA with Big Data

Nguyen Thuy Duong** · Sungmook Lim**†

**Dongguk Business School, Dongguk University-Seoul

Data envelopment analysis (DEA) is a tool for identifying best-practices when multiple performance metrics or measures are present for decision-making units (DMUs). As big data issue becomes an important area of supply chain and operations management, DEA is evolving into a data-oriented data science tool for benchmarking, performance evaluation, composite index construction and others. As the number of DMUs increases, the running-time to solve the standard DEA model sharply rises. Such situations are appearing more frequently in the era of big data. This issue could be an important challenge particularly when real-time data stream in at extremely high rates and the DEA analysis needs to be performed very quickly. Therefore, there exist practical needs for developing an efficient way of solving large-scale DEA problems. In this paper, we propose a practical approach for speeding up the DEA efficiency estimation process based on machine learning. In this approach, a sample of DMUs is selected from the population as a training data set, based on which a machine is trained to predict the efficiency scores of unselected or newly streamed-in DMUs. We also suggest a data augmentation technique to enhance the learning process under severe data class imbalance. The superior performance of the proposed approach over the conventional one in terms of efficiency prediction power as well as model computation time is shown through a series of computational experiments using randomly generated data.

Keyword : Data envelopment analysis, Large-scale, Big data, Machine learning, Support vector machine

* 이 논문은 교신저자에 대한 2021년도 동국대학교 연구년 지원에 의하여 이루어졌으며, 제1저자의 동국대학교 일반대학원 경영학과 석사학위 논문을 보완 및 확장한 것임.

† **Corresponding Author** : Dongguk Business School, Dongguk University, 30, Pildong-ro 1-gil, Jung-gu, Seoul, 04620, Korea. Tel: +82-2-2260-3814, E-mail: sungmook@dongguk.edu

Received: 3 December 2021, **Accepted**: 12 December 2021

1. 서 론

DEA(Data envelopment analysis)는 Charnes, Cooper & Rhode(1978)에 의해 개발된 이후 오랜 기간에 걸쳐 수많은 응용분야에서 폭넓게 성공적으로 활용되어 왔는데, 투입요소의 변환과정을 통해 산출요소를 생산하는 생산운영 프로세스의 효율성을 측정하는 전통적인 모형을 벗어나 다중 성능요인이 존재하는 상황에서 베스트 프랙티스 경계선(best-practice frontier)을 추정하고 이로부터의 거리를 기준으로 성과를 측정하는 일반화된 모형으로 점차 확장되어 왔다.

DEA를 투입-산출요소 간 변환과정의 생산 효율성을 분석하기 위한 방법론으로 협소하게 해석할 수도 있으나, 변수간 관계성을 분석하는 여타 방법론들과는 차별적으로 DEA가 가지는 모형의 단순성과 유연성으로 인해 다양한 응용과 해석이 가능하다. 특히, 다기준의사결정(multi criteria decision making, MCDM) 분야는 다양한 해법을 위한 유용한 토대로서 DEA가 활발하게 활용되어 온 분야 중 하나이다(Stewart, 1996). DEA에서의 DMU는 MCDM에서의 다수 대안으로 대응되고, DEA에서의 투입-산출요소는 MCDM에서의 다중 성능요인에 대응되며, DEA에서의 효율성 개념은 MCDM에서의 볼록 효율성 개념에 대응된다. DEA가 MCDM의 한 기법으로 사용되는 경우 DEA는 다중 요인 성능 측정 모형(multi-factor performance measurement model)으로 간주될 수 있다(Lim et al., 2014).

한편, DEA는 많은 데이터 포인트 중에서 경계선 이상치(frontier outliers)를 식별하는 도구로 이해될 수도 있다(Dulá & López, 2013). 경계선 이상치에 해당하는 데이터 포인트들은 그들이 속한 데이터 집합 내에서 개별적 또는 복합적으로 상한 또는 하한에 해당하는 특성값을 가진다는 점에서 특별한 주의를 요하는 데이터 포인트가 될 가능성이 높다. 이러한 용도로 DEA가 활용되는 경우 DEA는 기계학습이나 데이터과학에서 다루는 데이터 분류 모형의 하나로도 볼 수 있다.

DEA를 데이터 분석도구로서 인식하고 빅데이터와 연관 지어 연구하려는 시도가 최근 들어 이루어지고 있다. 이러한 시도는 두 가지 방향으로 이루어지고 있는데, 하나는 DEA를 통해 측정되는 효율성 척도를 예측변수(predictor variable) 또는 특징변수(feature variable)로 하여 평가대상 집합을 분류하거나 특정 목표변수(target variable) 또는 반응변수(response variable)을 예측하는 모형을 구성하는 연구이다. 이러한 유형의 연구 중 일부

는 사실 오래전부터 2단계 DEA 분석이라는 명칭으로 다수 수행되어 발표된 바 있으며, 조직의 효율성이 해당 조직의 또 다른 성과(예를 들어, 주가가격, 고객만족도 등)에 어떤 영향을 미치는지 또는 어떤 환경변수(예를 들어, 제도 변화, 조직 경영방법 등)에 의해 영향을 받는 지 분석하였다. 최근 들어서는 DEA를 확장하여 네트워크 구조를 가지는 DMU의 효율성을 각 단계별로 측정하는 네트워크 DEA 모형에 대한 연구가 활발해지면서, 단면적인 2단계 DEA 분석을 넘어 다차원적이고 복합적인 변수 관계 구조를 대상으로 하는 데이터 기반 애널리틱스 도구의 하나로서 DEA를 재구성하려는 시도도 나타나고 있다. 특히 Zhu(2020)는 네트워크 구조로 표현되는 빅데이터 내 은닉된 가치 있는 정보를 추출하기 위해 네트워크 DEA 모형을 활용할 수 있다고 하였다.

빅데이터와 관련한 DEA 모형에 관한 연구의 두 번째 방향은 DMU의 개수, 투입-산출요소의 개수 등이 기하급수적으로 많아져 최적화 모형의 규모가 대형화되는 경우 그 계산 과정을 효율화하는 방법에 관한 연구이다. 대형 DEA 모형의 풀이 과정을 개선하는 연구는 그간 종종 있어 왔는데, 최근 들어 빅데이터 관련 이슈가 부각되면서 빅데이터 하에서의 DEA 모형의 계산을 빠르게 하기 위한 연구가 다시 주목을 받고 있다(Charles et al., 2019; Khezrimotlagh et al., 2019; Shi et al., 2020). 평가대상 DMU의 개수가 아주 많거나 지속적으로 새로운 DMU가 대규모로 스트리밍(streaming)되는 경우에는 최적화 계산이 필요한 DEA 모형이 대형화되고 이를 풀기 위해 소요되는 계산시간은 최신의 최적화 소프트웨어를 사용한다고 하더라도 장시간이 소요될 수 있다. 여기서 장시간이라 함은 응용분야에 따라 상대적인 의미를 가지는데, 빅데이터 또는 스트리밍 데이터가 대규모로 발생되고 이를 빠르게 반영하여 DEA 평가가 이루어져야 하는 상황에서는 실시간에 가까운 계산시간이 요구될 수 있다. 예를 들어, 신용카드 결제 데이터를 분석하여 부정사용 패턴을 실시간으로 탐지하기 위해 경계선 이상치 식별도구로서 DEA를 활용할 수 있는데, 이때 신용카드 결제 데이터는 빅데이터, 스트리밍 데이터로 볼 수 있고 이에 대한 DEA 분석은 거의 실시간으로 이루어질 필요가 있다. 신용카드 결제 데이터는 상당한 규모의 빅데이터이고 실시간으로 스트리밍되어 누적되는 데이터의 양도 상당하므로, 이를 일반적인 DEA 모형 계산방식으로 풀어서는 해당 응용 요구사항에 맞는 신속한 평가가 이루어지기는 힘들다. 따라서 빅데이터 하에서의 DEA 평가가 빠르게 이루어질 수 있도록 하는 방법론의 연구가 필요해지며, 본 연구는 이러한 필요를 충족시키기 위한 목적을 가진다.

본 연구에서는 앞서 그 필요성을 설명한 빅데이터 DEA 모형의 효율적 계산을 위한 방법론으로 기계학습에 기반한 방법을 제안한다. 대략적으로 설명하면, 먼저 전체 DMU 집합 중 일정 비율만큼 표본추출하고 그 표본집합을 대상으로 DEA 모형을 풀어 효율성 점수를 계산한다. 이후 표본집합에 속한 DMU들을 훈련데이터(training data)로 하고 투입-산출요소를 특징변수로, 효율성 점수를 목표변수로 하는 기계학습 모형을 훈련시킨다. 이렇게 구성된 기계학습 모형을 이용하여 표본집합에 포함되지 않은 나머지 DMU들의 효율성 점수를 추정하고, 이후 스트리밍되어 신규로 추가되는 DMU들에 대해서도 그 효율성을 추정한다.

본 논문은 다음과 같이 구성된다. 먼저 DEA와 기계학습 등 기초적인 이론적 배경을 설명하고, 이어서 제안하는 방법론을 제시한다. 이후 제안 방법론의 유효성을 보이기 위해 실증 실험을 진행한 결과를 보인다. 마지막으로 토의와 함께 결론을 제시한다.

2. 이론적 배경

2.1 DEA

DEA는 다수의 성과척도(performance measures)를 가지는 의사결정단위(Decision Making Unit, DMU)의 상대적 성과지표(performance index)를 도출하기 위한 수리 최적화 기반의 비모수적 모형이다. DEA는 원래 투입요소의 변환과정을 통해 산출요소를 생산하는 생산운영 프로세스의 효율성을 측정하는 모형으로 개발되었다. DEA를 통한 효율성 평가대상이 되는 각 의사결정단위의 운영 프로세스는 기업의 경우를 예로 든다면 노동, 자본 등의 투입요소를 제품 또는 서비스 등의 산출요소로 변환하는 과정으로 이해할 수 있다. 이러한 운영 프로세스가 얼마나 효율적인지를 측정하고 평가하려는 연구는 생산경제학의 전통적인 주제이며, 가용 자원 조건 하에서 다른 산출요소 수준을 악화시키지 않으면서 어느 한 산출요소 수준을 개선시킬 수 있는지 여부, 즉 Pareto(1927) 조건의 개념을 효율적 운영 프로세스의 요건으로 정의한 Koopmans(1951)가 대표적인 초기 연구이다. 이후 Pareto와 Koopmans의 효율성 개념을 실제 데이터에 적용하여 상대적 효율성을 측정할 수 있는 방법을 처음으로 제시한 Farrell(1957)의 연구를 토대로 그 한계점을 극복하고 Pareto와 Koopmans의 효율성을 완전한 형태로 구현할 수 있는 선형계획법 기반의 방법이 제안

되기에 이르는데, 이 방법이 Charnes, Cooper & Rhode (1978)에 의해 개발된 DEA이다.

DEA 모형은 생산가능집합을 구성하기 위한 공리적 가정들을 어떻게 설정하는지, 그리고 효율적 경계선까지의 거리를 어떻게 측정하는지에 따라 다양한 모형으로 나누어진다. 먼저, 생산가능집합을 구성하기 위한 가정들로는 자유가처분성, 볼록성, 규모수익성 등이 있는데, 그 선택 조합에 따라 FDH(free disposal hull) 모형, CRS(constant returns-to-scale) 모형, VRS(variable returns-to-scale) 모형 등이 도출된다. 또한, 효율적 경계선까지의 거리를 측정하는 함수의 형태에 따라 방사형(radial) 모형, SBM(slacks-based measure) 모형, 방향거리함수(directional distance function) 모형 등이 도출된다.

본 절에서는 다양한 DEA 모형 중 대표적이면서 본 연구에서 사용된 BCC 모형과 CCR 모형에 대해 소개한다. 두 모형 모두 거리함수로 투입지향 또는 산출지향 방사형 모형을 취하고, 생산가능집합에 대한 공리로 자유가처분성과 볼록성을 가정한다는 점에서 동일하다. 그러나 생산가능집합의 규모수익성 가정에서 차이가 있는데, BCC 모형은 VRS 모형이고 CCR 모형은 CRS 모형이다. 이들 모형은 다음과 같이 정형화된다.

n 개의 DMU가 존재하고, 각 DMU는 동질적인 투입-산출구조를 가지면서 m 개의 투입요소를 이용해 s 개의 산출요소를 생산한다고 하자. DMU $j(j \in \{1, 2, \dots, n\})$ 가 가지는 투입요소의 양을 $x_j = (x_{1j}, \dots, x_{mj})^T \in R_+^m$, 산출요소의 양을 $y_j = (y_{1j}, \dots, y_{sj})^T \in R_+^s$ 라고 할 때, 투입지향 BCC 모형에서는 DMU j 의 효율성 점수를 다음의 선형 최적화 모형을 풀어 구한다.

$$\begin{aligned} \min \quad & \theta \\ \text{s.t.} \quad & (\theta x_j, y_j) \in P. \end{aligned} \quad (1)$$

여기서 $P = \{(x, y) \mid X\lambda \leq x, Y\lambda \geq y, e^T \lambda = 1, \lambda \geq 0\}$ 는 생산가능집합을 의미하고, 모든 DMU들의 투입-산출요소 값을 행렬의 형태로 정리한 X 와 Y 는 각각 $X = (x_{ij}) \in R^{m \times n} (i = 1, \dots, m)$, $Y = (y_{kj}) \in R^{s \times n} (k = 1, \dots, s)$ 로 정의된다. P 의 정의에서 $e^T \lambda = 1$ 이라는 제약조건을 제외시키면 모형 (1)은 CCR 모형이 된다. 모형 (1)의 최적 목적함수 값 θ^* 가 평가대상 DMU의 효율성 점수가 된다.

한편, 방사형 거리함수와는 달리 방향거리함수를 사용하는 DEA 모형은 다음과 같이 정형화 된다.

$$\begin{aligned} \max \quad & \delta \\ \text{s.t.} \quad & (x_j - \delta d_x, y_j + \delta d_y) \in P. \end{aligned} \quad (2)$$

여기서 $d = (d_x, d_y)$, $d_x \in R_m^+$, $d_y \in R_s^+$ 로 투입요소 및 산출요소의 개선방향을 나타낸다. 최적 목적함수 값 δ^* 은 평가대상 DMU의 투입-산출요소의 값을 개선방향 d 를 따라 얼마나 이동하면 효율적 경계선에 도달하는지를 나타내며, 그 값이 0이면 해당 DMU가 효율적이라는 것을 의미하고 그 값이 커질수록 비효율성이 커짐을 뜻한다. 다만, 방사형 모형과는 달리 δ^* 의 값은 0과 1 사이에 있지는 않고, 개선방향의 설정, 투입-산출요소의 척도 등에 따라 상한값은 달라진다.

2.2 대규모 DEA 계산을 위한 방법

통상적인 DEA 평가 과정에서는 각 개별 DMU별로 한번씩, 총 n 번의 선형계획법 모형을 풀어야 한다. 하나의 DMU에 대한 DEA 모형 (1)은 변수의 비음조건 이외 $(m+s+1)$ 개의 제약조건과 n 개의 변수를 가지는 선형계획모형인데, 이러한 선형계획모형을 평가대상 DMU를 바꿔가며 n 번 풀어야 한다. DMU의 개수가 커지면 DEA 모형의 계산량은 급속하게 증가하는데, 이러한 문제는 빅데이터 기반의 데이터 분석이 요구되는 상황에서 더욱 자주 나타나게 된다.

대규모 DEA 평가에 소요되는 계산량 이슈는 고빈도로 데이터가 실시간 스트리밍되고 DEA 분석이 그에 맞추어 아주 신속하게 이루어져야 하는 상황에서는 특히 중요한 도전 과제가 될 수 있다. 그러므로 대규모 DEA 문제를 효율적으로 풀어내는 방법의 개발은 그 실용적 필요성이 크다고 할 수 있다.

대규모 DEA 문제를 효율적으로 푸는 방법에 관한 연구는 많이 존재하는데, 대부분 문제의 크기를 축소시키거나 전체 문제를 몇 개의 부분 문제로 분할하여 푸는 방법을 통해 계산량을 줄이는데 초점을 두고 있다. Dulá and Helgason(1996)은 유한한 DMU 집합에 대해 선형계획모형의 크기가 효율적 경계선 상의 극점에 위치한 DMU의 개수보다 크지 않도록 줄이는 계산절차를 개발하였고, Barr and Durchholz(1997)는 대규모 DEA 문제를 푸는 계산시간을 줄이기 위해 계층적 분해 절차를 제안하였다. 한편, Dulá(2008)는 DMU의 개수, 투입-산출요소의 개수, 그리고 효율적 DMU의 비율 등 세 가지 파라미터가 DEA 문제의 계산시간에 미치는 영향을 분석하였다. Dulá(2011)는 모든 효율적 DMU를 찾는데 소요되는 계산시간을 크게 줄일 수 있는 'build-hull'이라는 절차를 개발하였고, Zhu et al.(2018)은 빅데이터 환경에서 DMU들을 여러 부분집합으로 분할함으로써 계산과

정을 가속화하는 새로운 알고리즘을 개발하였다. Chen & Lai(2017)는 대규모 데이터에 대한 방사형 효율성을 소규모 선형계획모형을 풀어 계산하는 알고리즘을 개발하였는데, 해의 품질을 유지하면서 매번 100개의 데이터 포인트 이하로 개별 선형계획모형의 크기를 통제하는 방법을 취하였다. Dulá & López(2013)는 데이터가 고빈도로 스트리밍되어 유입될 때 DEA 모형 계산을 빠르게 하는 방법을 연구하였고, 스트리밍 데이터를 처리하는 이론적 토대와 함께 얼마나 신속하게 데이터가 처리될 수 있는지에 대해 논의하였다. Dulá & Thrall(2001)은 DEA 계산 시간을 단축하고 응용에서의 유연성을 증가시키는 프레임워크를 제안하였는데, 주어진 모형을 표현하기 위해 필요한 최소한의 데이터 부분집합을 식별하는 방식에 기초하였다. Chen & Cho(2009)는 효율성 점수를 결정하는데 핵심적인 역할을 하는 몇 개의 유사한 DMU를 식별하는 방식으로 DEA 계산절차를 가속화하는 방안을 제안하였다. Sueyoshi & Chang(1989)는 DEA 모형의 특성을 고려하여 선형계획문제를 푸는 심플렉스 방법 기반의 특수 알고리즘을 개발하여 DEA 계산 속도를 개선시켰다.

3. 제안 프레임워크

본 연구에서는 빅데이터 또는 스트리밍 데이터 하에서 DEA를 통한 효율성 추정을 실시하는데 기계학습 방법을 적용하는 프레임워크를 제안한다. 기본적인 아이디어는 우선 전체 DMU들 중 일부를 훈련데이터로 표본추출한 후 이를 이용해 투입-산출요소를 특징변수로 하고 효율성 점수를 목표변수하여 기계학습 모형을 훈련시킨다. 다음으로 훈련된 기계학습 모형을 이용하여 나머지 DMU들로 구성되는 평가데이터(test data) 또는 새롭게 스트리밍되어 추가된 DMU를 대상으로 효율성 점수를 추정한다. 한편, 학습과정이 좀 더 빠르게 진행될 수 있도록 하는 데이터 증강법(data augmentation)을 제안하는데, 훈련데이터 DMU들을 효율적 경계선으로 완전히 또는 부분적으로 투영하여 효율적 경계선 근방의 데이터를 인공적으로 다수 생성시킨다.

본 연구에서 제안하는 상기 접근법을 통해 한편으로는 빅데이터 또는 스트리밍 데이터 하에서의 DEA 효율성 추정의 과도한 계산량 부담을 실질적으로 줄이면서 그에 따른 정확도 저하를 최소한으로 할 수 있을 것으로 기대한다. 제안하는 방법의 수행절차를 요약하면 다음과 같다.

- ▶ 단계 1: (훈련데이터 생성) 전체 DMU 모집단으로부터 일정 비율(ρ) 크기의 표본을 훈련데이터로 추출한 후, 해당 훈련데이터에 대해 DEA 효율성 점수를 계산한다.
- ▶ 단계 2: (데이터 증강) 비효율적인 DMU들을 효율적 경계선으로 완전히(투영 후 효율성 점수가 1이 됨) 그리고 부분적으로(투영 후 효율성 점수가 0.9가 됨) 투영하여 가상의 DMU들을 생성하고, 이를 포함시켜 원 훈련데이터를 증강시킨다. 증강된 DMU들의 효율성 점수는 별도의 계산이 불필요하다.
- ▶ 단계 3: (학습) 단계 1과 단계 2에서 계산된 DEA 효율성 점수를 목표변수로 하고 각 DMU별 투입-산출요소 값을 특징변수로 하여, 증강된 훈련데이터로 기계학습 모델을 훈련시킨다.
- ▶ 단계 4: (예측) 훈련된 기계학습 모델을 이용하여 전체 DMU 모집단 중 훈련데이터에 포함되지 않았거나 향후 단기간에 걸쳐 새롭게 스트리밍되어 유입되는 DMU들의 효율성 점수를 추정한다.

상기 절차를 도식화하면 <Fig. 1>과 같다.

단계 1에서 결정하는 비율 ρ 는 풀어야 하는 DEA 모형의 크기를 결정하는 동시에 기계학습 모형을 위한 훈련데이터의 크기를 결정한다. 즉, 비율 ρ 의 값을 크게 할수록 DEA 모형의 크기는 커져 계산량 부담은 커지지만 더 큰 훈련데이터로 학습이 이루어지므로 학습모형의 예측 성능은 높아질 것을 기대할 수 있다. 다만, 훈련데이터의 크기가 커질수록 기계학습 과정에 소요되는 계산시간은 늘어난다는 점은 고려해야 한다. 한편, 비율 ρ 의 값을 작게 하면 DEA 모형의 크기는 작아져 계산량 부담은 낮

아지만 훈련데이터도 함께 작아져 학습모형의 예측 성능이 저하될 수 있다. 다만, 훈련데이터의 크기가 작아지면 학습과정의 소요시간은 짧아진다. 이와 같이, 비율 ρ 는 계산량 부담과 학습모형 성능 간의 트레이드오프를 결정한다고 볼 수 있는데, 종합하자면 ρ 의 값을 크게 하면 DEA 모형 계산 및 기계학습 시간이 늘어나지만 구축된 학습모형의 예측력은 향상되고, ρ 의 값을 작게 하면 DEA 모형 계산 및 기계학습 시간은 줄어들지만 학습모형의 예측력은 저하되는 경향이 있다.

단계 2에서 이루어지는 데이터 증강은 학습과정의 성능을 큰 비용 없이 개선시킬 수 있는 방법이다. 데이터 증강은 훈련데이터의 양이 충분하지 않거나 데이터의 불균형 등이 있을 때 이를 해소하기 위해 추가적인 데이터를 일정한 방식에 따라 인위적으로 생성하여 훈련데이터를 보강하는 방법으로서, 추가적인 데이터를 확보하는 과정에 소요되는 비용이 크지 않다면 효과적으로 학습과정의 성능을 향상시킬 수 있는 방법이 될 수 있다. 본 연구에서 다루고 있는 DEA 모형을 통한 효율성 평가의 경우 비효율적인 DMU와 효율적인 DMU 간에 데이터 불균형이 심할 수 있고, 특히 대규모 DEA 모형에서는 DMU의 개수가 투입-산출요소 개수보다 훨씬 많은데 이 경우에는 DEA 모형의 특성상 변별력(discriminating power)이 올라가 효율적인 DMU의 개수가 상대적으로 희박해지는 경향이 있다. 따라서 본 연구에서 제안하는 학습과정의 성능을 보다 높이기 위해서는 데이터 증강이 요구된다.

본 연구에서는 DEA 모형의 특성을 활용한 데이터 증강 방법을 제안한다. 즉, DEA 평가에서 비효율적이라고 판별되는 DMU는 효율적 경계선으로의 투영을 통해 효

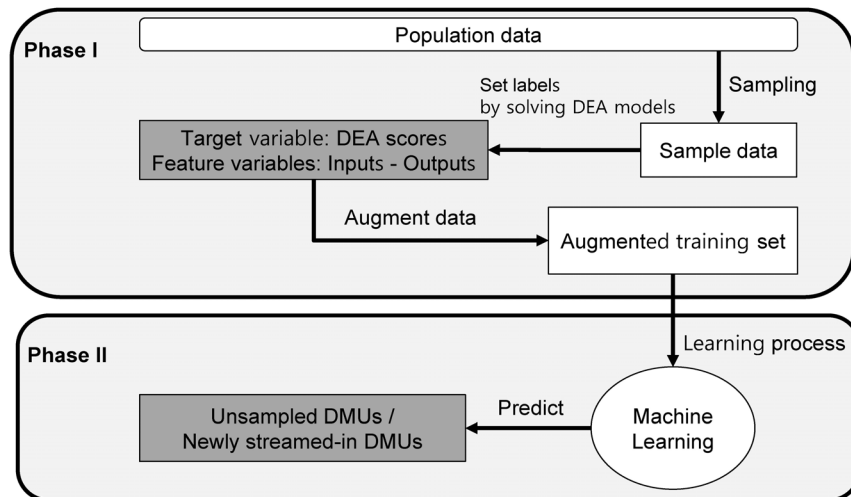


Fig. 1 Proposed Framework

율적인 DMU로 변환시킬 수 있다는 점을 활용하여, 비효율적인 DMU 데이터로부터 효율적인 DMU 데이터를 인위적으로 생성해낼 수 있다. 앞서 살펴본 DEA 모형 (1)을 풀어 얻어진 DMU j 의 효율성 점수를 θ^* 라고 한다면, 투영점 $(\theta^* x_j, y_j)$ 는 효율적 경계선 상에 존재하므로 해당 투영점을 효율적 DMU 데이터로 추가할 수 있다. 이와 더불어, 부분적 투영을 통해 효율적이지는 않으나 효율성 점수가 1에 가까운 DMU를 생성할 수도 있다. 예를 들어, 투영점의 투입요소 값을 10% 증가시킨 점인 $(\theta^* x_j/1.1, y_j)$ 은 효율성 점수가 1/1.1인 비효율적인 DMU가 되며, 이를 추가하여 훈련데이터를 증강시킬 수 있다. 부분적 투영을 통한 데이터 증강은 효율적 DMU와 비효율적 DMU를 이진 분류하는 학습모형에 분류 난이도가 높은 훈련데이터를 충분히 제공함으로써 학습 성능을 높이기 위한 목적을 가진다. 데이터 증강을 얼마나 할지는 학습 성능과 계산량의 트레이드오프 문제가 된다. 즉, 더 많은 비효율적 DMU를 투영하여 증강 데이터를 생성할수록 학습 성능의 개선은 기대할 수 있지만 학습에 소요되는 계산량은 늘어나고, 그 반대로 성립한다.

데이터 증강 과정을 예시로 적용하면 <Fig. 2>와 같다. <Fig. 2>의 상단 왼쪽 그림은 1개의 투입요소와 1개의

산출요소를 가지는 71개의 DMU를 임의로 생성하고 이를 효율적 경계선과 함께 2차원 좌표평면 상에 표시하고 있다. 이 예시에서는 효율적 DMU의 개수는 5개이고 비효율적 DMU는 66개로서 데이터 분류별 불균형이 심하다. 이러한 문제를 해소하여 학습 성능을 향상시키기 위해 모든 비효율적 DMU를 효율적 경계선으로 완전히 투영하여 데이터를 증강한 결과는 <Fig. 2> 상단 오른쪽 그림과 같다. 더불어, 비효율적 DMU를 부분적으로 효율적 경계선 방향으로 투영하여 데이터를 추가로 증강한 결과는 <Fig. 2> 하단의 그림과 같다.

단계 3에서 사용할 기계학습 모형으로는 응용 목적에 따라 다양한 선택이 가능하다. 효율적 DMU와 비효율적 DMU를 단순히 이진 분류하기 위한 목적이라면 분류모형(classification machine)에 해당하는 학습모형을 선택하고, 효율성 점수 자체를 추정하기 위한 목적이라면 회귀모형(regression machine)에 해당하는 학습모형을 선택하여야 한다. 어떤 기계학습 모형이 DEA 효율적 경계선을 식별하고 그로부터의 거리로 측정되는 효율성 점수를 잘 예측하는지에 대해 본 연구에서는 실증 실험을 실시하는데, 그 결과는 다음 절에서 제시한다.

단계 4에서는 단계 3을 통해 훈련된 기계학습 모형을

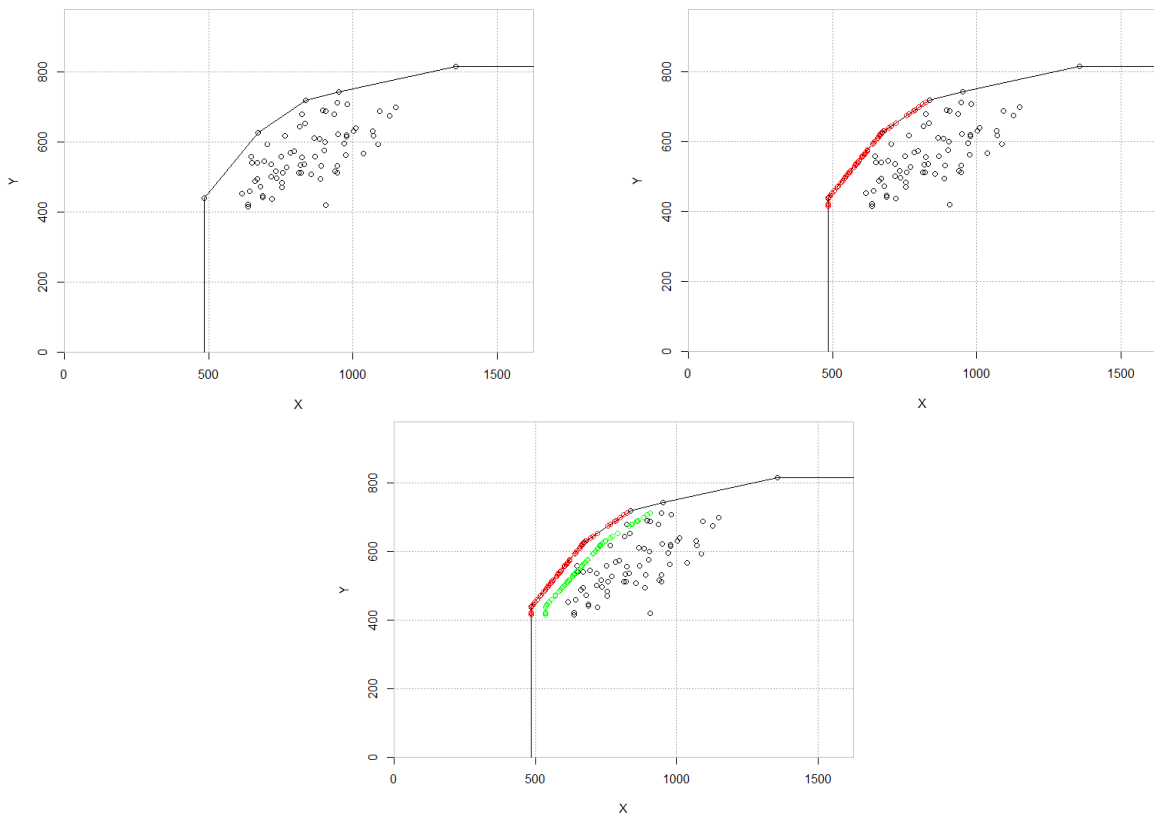


Fig. 2 Data Augmentation Example

이용하여 전체 DMU 모집단 중 훈련데이터에 포함되지 않았거나 향후 단기간에 걸쳐 새롭게 스트리밍되어 유입되는 DMU들의 효율성 점수를 추정한다. 이 과정은 계산량도 미미하고 단순하게 수행될 수는 있으나 논의가 필요한 부분이 있는데 이는 다음과 같다. 훈련데이터에 포함되지 않은 DMU들은 훈련데이터에 포함된 DMU와 동일한 시점과 환경에서 생성된 데이터이므로 동일한 효율적 경계선을 공유한다고 볼 수 있고, 따라서 해당 효율적 경계선의 특성을 훈련데이터만으로 학습한 후 훈련데이터에 포함되지 않은 DMU의 효율성 점수를 예측하는 과정의 논리는 타당하다고 볼 수 있다. 예측에서 발생하는 오차는 표본추출에 따른 오차와 기계학습 모형의 용량 문제 또는 과소적합 및 과적합에 기인한다. 반면, 새롭게 스트리밍되어 유입되는 DMU의 경우에는 기존에 존재했던 DMU와는 발생시점 및 환경 등에 있어 이질적일 수 있으므로 기존 DMU와 효율적 경계선을 공유하지 않을 수 있다. 이 경우 훈련데이터에 의해 학습된 효율적 경계선을 기준으로 새로 유입된 DMU의 효율성을 예측하는 것은 타당하지 않을 수 있다. 그러므로 본 연구에서 제안하는 방법을 통해 스트리밍 데이터에 대한 효율성 예측을 수행하는 것이 타당성을 가지려면 분석 대상 기간 내에 효율적 경계선의 구조적 변화가 일어나지 않아야 한다는 가정이 필요하다. 따라서 본 연구의 방법은 단기간에 걸쳐 새롭게 스트리밍되어 유입되는 데이터에 대해서만 적용되는 것이 바람직하며, 일정 기간이 경과한 이후에는 변화된 훈련데이터를 대상으로 기계학습 모형 구축을 다시 할 필요가 있다.

4. 실증 실험

4.1 실험 설계

본 연구에서 제안한 접근법의 타당성을 보이기 위해 실증 실험을 실시하였다. 실증 실험을 위한 코드는 R ver. 4.0.3과 RStudio ver. 1.4.1103을 기반으로 DEA 모형 및 기계학습을 위한 패키지를 설치, 활용하여 작성하였다.

실험 설계를 위해 결정해야 하는 사항은 ① DEA 모형의 유형, ② 투입-산출요소 및 DMU의 수, ③ 기계학습 모형의 유형, ④ 훈련-평가데이터의 분할비율(ρ), ⑤ 학습 성능평가 척도 등이다. 우선, DEA 모형의 유형을 결정하는 구조적 요인은 효율적 경계선의 형태와 거리함수의 유형이므로, 이를 기준으로 여러 유형의 DEA 모형을 실험 대상으로 포함한다. 효율적 경계선의 형태는 규모수

익성에 대한 가정에 따라 결정되는데, 본 논문에서는 가장 자주 사용되는 두 가지 가정인 불변규모수익성(CRS)과 가변규모수익성(VRS)에 대해 실험한다. 한편, 효율성 점수를 계산하기 위한 평가대상 DMU와 효율적 경계선 간의 거리함수로는 방사거리(radial distance)함수와 방향거리(directional distance)함수를 두 가지 대표적인 유형으로 선택하여 실험한다. 거리함수의 유형으로 투입지향 모형과 산출지향 모형을 구분할 수도 있는데, 본 실험의 형태상 두 모형간 차이가 없으므로, 투입지향 모형을 택한다. 상기와 같은 방식으로 구성되는 DEA 모형의 유형은 규모수익성 가정과 거리함수 유형의 조합에 따라 총 4가지가 실험 대상이 된다.

투입-산출요소 및 DMU의 수는 DEA 모형의 규모와 이에 따른 계산량을 결정한다. 본 실험에서 DMU의 개수는 1,000개, 3,000개, 5,000개, 10,000개 등 4가지로 하고, 투입요소 및 산출요소의 수는 (2, 2), (3, 3), (5, 5) 등 3가지 조합으로 실험한다. 각 DMU의 투입요소 및 산출요소의 값은 일정 규칙에 따라 랜덤하게 생성하며 상세한 내용은 다음 절에서 설명한다. DEA 모형의 계산은 Benchmarking 패키지 ver. 0.29를 사용하였다.

기계학습 모형의 유형은 분류모형과 회귀모형으로 각각 실험하였는데, 분류모형으로는 SVM, 랜덤포레스트, 로지스틱회귀분석 모형을, 회귀모형으로는 SVR, 랜덤포레스트, 선형회귀분석을 적용한다. 기계학습 모형에서의 예측변수는 투입요소 및 산출요소 값이고 목표변수는 특정 DEA 모형에 의해 계산되는 효율성 점수로 한다. 분류모형에서는 효율성 점수가 1인 경우(효율적)와 1 미만인 경우(비효율적)로 구분하는 이진 목표변수를 취하며, 회귀모형에서는 효율성 점수의 값 그 자체를 연속 목표변수로 취한다. SVM(SVR)은 e1071 패키지 ver. 1.7-8, 랜덤포레스트는 randomForest 패키지 ver. 4.6-14을 사용하였고, 선형회귀분석과 로지스틱회귀분석은 stats 패키지 ver. 4.0.3를 사용하였다. 모든 모형에서 해당 패키지의 디폴트 파라미터 값을 유지하고 별도의 튜닝은 하지 않았다.

훈련데이터와 평가데이터의 분할비율은 50%-50%와 70%-30%로 각각 설정하여 비교 실험한다. 데이터 분할은 효율적 DMU와 비효율적 DMU 간의 비율이 훈련데이터와 평가데이터에서 동일하게 유지되도록 하였으며, 이를 위해 R(ver. 4.0.3)의 caret 패키지 ver. 6.0-88에 포함된 createDataPartition 함수를 활용하였다. 훈련데이터의 증가는 훈련데이터 내 비효율적 DMU의 10%를 임의로 선택해 효율적 경계선으로 완전히 투영하고, 동시에 투영점의 투입요소 값을 10% 증가시켜 효율적 경계선

에 가까운 데이터로 부분적 투영을 하는 방식으로 실시한다.

평가데이터를 대상으로 하는 학습모형 예측성능 평가척도로는 분류모형의 경우 혼동행렬(Confusion matrix)을 기초로 통상적으로 사용되는 척도를 사용하는데, 정확도(Accuracy)를 비롯해 민감도(Sensitivity)와 특이도(Specificity)의 산술평균인 균형정확도(Balanced accuracy), 정밀도(Precision)와 재현율(Recall)의 조화 평균인 F1 점수 등을 확인한다. F1 점수와 균형정확도는 클래스간 불균형이 심한 경우에 예측성능을 효과적으로 나타낼 수 있다. 회귀모형의 경우 결정계수 R^2 을 비롯해 Spearman 순위상관계수, 예측오차의 크기를 보기 위한 평균제곱근오차(Root Mean Squared Error, RMSE) 등을 확인한다. 평가데이터의 효율성 점수는 데이터 분할 전 모집단 DMU를 대상으로 하는 DEA 모형을 풀어 계산된 점수이며, 이 점수가 학습모형의 예측결과와 얼마나 부합하는 지로서 예측성능을 측정한다. 한편, 여러 기계학습 모형 간의 성능차이를 통계적으로 검정하기 위해 윌콕슨 대응표본 부호순위 검정(Wilcoxon paired sample signed rank test)을 적용하여 확인한다.

4.2 데이터 생성

실험에서 사용할 DMU 투입-산출요소 데이터는 다변

량 정규분포에 기초하여 랜덤하게 생성한다. 우선, 각 투입-산출 요소별 평균값은 100에서 1000사이의 일양분포로부터 추출하고, 각 요소별 표준편차는 10%에서 30% 사이의 일양분포로부터 개별적으로 추출된 값을 해당 평균값에 곱하는 방식으로 정한다. 한편, 투입요소와 산출요소들 간의 상관관계는 다음과 같이 고려하였다. 먼저, 투입요소 간의 독립성을 가정하여 상관계수를 0으로 하고, 산출요소 간의 상관계수도 0으로 한다. 그러나 투입요소와 산출요소 간에는 투입에 따른 산출 증가라는 일반적인 관계성을 반영하기 위해 투입-산출요소 간 양의 상관관계를 가정한다. 각 투입-산출요소간 상관계수의 크기는 전체적인 상관계수 행렬이 대칭양정치(symmetric positive definite)가 되는 조건 하에서 랜덤하게 생성한다. 이 과정에서 랜덤한 상관계수 행렬의 생성은 R(ver. 4.0.3)의 randcorr 패키지 ver. 1.0에 포함된 randcorr 함수를 활용하였고, 다변량 정규분포를 따르는 난수의 생성은 MASS 패키지 ver. 7.3-53에 포함된 mvrnorm 함수를 활용하였다.

〈Fig. 3〉은 투입-산출요소의 개수가 각각 3개인 700개의 DMU를 랜덤하게 생성한 데이터의 분포를 예시로서 보여주며, 이는 본 실험에서 실제 사용되었다. 참고로 그림에서 $x.1 \sim x.3$, $y.1 \sim y.3$ 은 각각 투입요소와 산출요소를 의미하고, 우삼각행렬 상자 내의 숫자는 교차하는 변수간 상관계수를 나타낸다.

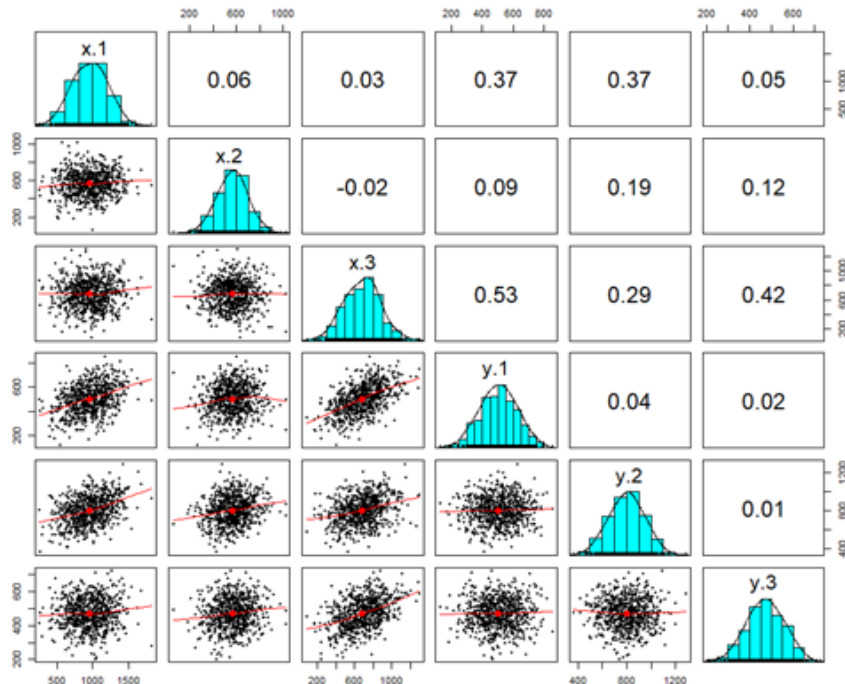


Fig. 3 Randomly Generated Data Example with 3 inputs, 3 outputs and 700 DMUs

4.3 학습 성능 비교

DEA 모형의 유형, 투입-산출요소 및 DMU의 수, 기계학습 모형의 유형별로 구성되는 다양한 조합별로 학습 성능평가 척도의 계산결과는 <Table 1>과 <Table 2>와 같다. 단, 훈련데이터와 평가데이터의 분할비율은 50%-50%로 설정하였다. 참고로 실험 결과 표에서

DEA 모형 구분으로 Rad는 방사형 모형을, Dir는 방향거리함수 모형을 뜻하고, Density는 전체 DMU들 중 효율적 DMU의 비율을 의미한다. 또한, RF는 랜덤포레스트, LoR은 로지스틱회귀모형, LR은 선형회귀모형을 뜻한다.

표에 제시된 실험 결과는 각각의 실험 파라미터 조합에 대해 30번의 반복 시행을 한 결과를 평균한 수치를 제시하고 있다. 실험에 활용한 분류모형 3가지와 회귀모

Table 1. Classification Performance Comparison between Different Settings ($\rho = 0.5$)

Size (n, m, s)	DEA Model		Density	Accuracy			F1			Balanced accuracy		
				SVM	RF	LoR	SVM	RF	LoR	SVM	RF	LoR
(1000, 2, 2)	Rad	VRS	2.7%	0.9736	0.9684	0.9685	0.5276	0.5330	0.4473	0.7755	0.8147	0.7298
		CRS	0.8%	0.9917	0.9891	0.9877	0.5202	0.4682	0.4296	0.7793	0.7850	0.7903
	Dir	VRS	2.9%	0.9748	0.9664	0.9683	0.4943	0.5188	0.4118	0.7366	0.8135	0.6908
		CRS	1.0%	0.9888	0.9832	0.9862	0.4610	0.4277	0.4633	0.6785	0.7790	0.7709
(1000, 3, 3)	Rad	VRS	8.6%	0.9364	0.9179	0.9030	0.6729	0.5902	0.4572	0.8612	0.8138	0.7113
		CRS	3.8%	0.9642	0.9586	0.9553	0.5263	0.4712	0.4653	0.7718	0.7414	0.7515
	Dir	VRS	8.6%	0.9369	0.9236	0.9068	0.6584	0.5981	0.4603	0.8479	0.8117	0.7110
		CRS	3.8%	0.9678	0.9635	0.9586	0.4908	0.4957	0.4319	0.7230	0.7394	0.7125
(1000, 5, 5)	Rad	VRS	31.6%	0.8353	0.7860	0.7442	0.7837	0.7227	0.6407	0.8632	0.8093	0.7346
		CRS	24.4%	0.8596	0.8189	0.7748	0.7556	0.6709	0.5872	0.8725	0.7984	0.7343
	Dir	VRS	34.5%	0.8372	0.7680	0.7127	0.7963	0.7178	0.6223	0.8543	0.7817	0.6970
		CRS	23.5%	0.8655	0.8229	0.7858	0.7522	0.6565	0.5736	0.8665	0.7866	0.7253
(3000, 2, 2)	Rad	VRS	1.0%	0.9857	0.9845	0.9843	0.4961	0.4986	0.3522	0.8316	0.8651	0.7013
		CRS	0.4%	0.9950	0.9945	0.9944	0.5218	0.4598	0.4726	0.8111	0.7907	0.7899
	Dir	VRS	1.1%	0.9893	0.9873	0.9875	0.5749	0.5508	0.3561	0.8336	0.8538	0.6618
		CRS	0.3%	0.9961	0.9952	0.9953	0.5137	0.4923	0.4427	0.7968	0.8222	0.7671
(3000, 3, 3)	Rad	VRS	3.8%	0.9643	0.9600	0.9541	0.6141	0.5256	0.3890	0.8658	0.7861	0.6850
		CRS	1.7%	0.9827	0.9834	0.9787	0.5510	0.5074	0.4194	0.8315	0.7577	0.7249
	Dir	VRS	3.6%	0.9707	0.9673	0.9608	0.6652	0.5865	0.4048	0.8936	0.8159	0.6796
		CRS	1.6%	0.9853	0.9841	0.9814	0.5449	0.4781	0.4121	0.7958	0.7301	0.7108
(3000, 5, 5)	Rad	VRS	20.2%	0.8901	0.8557	0.8010	0.7732	0.6923	0.5337	0.9044	0.8320	0.7081
		CRS	12.2%	0.9198	0.9023	0.8769	0.7267	0.6203	0.5152	0.9076	0.7946	0.7309
	Dir	VRS	19.8%	0.8920	0.8564	0.8090	0.7718	0.6854	0.5163	0.9026	0.8275	0.6972
		CRS	13.2%	0.9191	0.8987	0.8676	0.7374	0.6296	0.4932	0.9014	0.7963	0.7073
(5000, 2, 2)	Rad	VRS	0.7%	0.9918	0.9908	0.9910	0.5346	0.4891	0.3373	0.8377	0.8241	0.6684
		CRS	0.2%	0.9969	0.9968	0.9968	0.5338	0.4581	0.4442	0.8490	0.7815	0.7449
	Dir	VRS	0.6%	0.9937	0.9929	0.9931	0.5796	0.5584	0.3759	0.8639	0.8773	0.6845
		CRS	0.2%	0.9973	0.9968	0.9967	0.4745	0.4142	0.4418	0.7706	0.7654	0.7508
(5000, 3, 3)	Rad	VRS	2.7%	0.9742	0.9734	0.9660	0.6256	0.5506	0.3691	0.8886	0.7942	0.6748
		CRS	1.1%	0.9889	0.9896	0.9866	0.5710	0.4743	0.3989	0.8587	0.7276	0.7085
	Dir	VRS	2.9%	0.9757	0.9712	0.9658	0.6505	0.5489	0.3489	0.8880	0.7932	0.6518
		CRS	1.2%	0.9887	0.9874	0.9861	0.5683	0.4602	0.3891	0.8198	0.7275	0.6897
(5000, 5, 5)	Rad	VRS	15.6%	0.9128	0.8866	0.8418	0.7657	0.6739	0.4930	0.9193	0.8302	0.6977
		CRS	8.5%	0.9432	0.9307	0.9138	0.7239	0.6020	0.4862	0.9198	0.7900	0.7167
	Dir	VRS	15.6%	0.9116	0.8853	0.8494	0.7665	0.6744	0.4859	0.9194	0.8326	0.6879
		CRS	8.7%	0.9422	0.9316	0.9105	0.7242	0.6030	0.4681	0.9163	0.7830	0.7027

형 3가지의 기계학습 모형 간에 여러 가지 예측 성능척도 기준의 비교우위를 보기 위해 30번 반복 시행 결과를 토대로 월록순 대응표본 부호순위 검정을 실시하고, 그 결과 유의한 우위를 보이는 항목은 굵은 숫자와 밑줄로 표시하였다.

실험 결과 거의 모든 경우에서 기계학습에 의한 예측 성능이 우수하게 나타나고 있다는 것을 확인할 수 있다.

분류모형의 경우 데이터 클래스 불균형이 심할 때 예측 성능을 잘 표현하는 F1 점수와 균형정확도가 모두 양호하게 나타나고 있고, 데이터 규모가 커질수록 성능이 더 향상되는 것을 알 수 있다. 회귀모형의 경우에도 결정계수가 거의 90% 이상 나타나고 있고, 순위상관계수는 95% 이상을 보이고 있다. RMSE로 측정되는 예측오차 또한 낮은 수준이다. 특히 분류모형 보다는 회귀모형에서 더

Table 2. Regression Performance Comparison between Different Settings ($\rho = 0.5$)

Size (n, m, s)	DEA Model		Den-sity	R^2			Spearman correlation			RMSE		
				SVR	RF	LR	SVR	RF	LR	SVR	RF	LR
(1000, 2, 2)	Rad	VRS	2.7%	<u>0.9141</u>	0.8162	0.7302	<u>0.9610</u>	0.9139	0.8760	<u>0.0431</u>	0.0596	0.0658
		CRS	0.8%	<u>0.9085</u>	0.8284	0.8003	<u>0.9541</u>	0.9141	0.9055	<u>0.0447</u>	0.0606	0.0642
	Dir	VRS	2.9%	<u>0.9296</u>	0.8258	0.7684	<u>0.9682</u>	0.9130	0.8752	<u>9.0228</u>	13.5404	13.6483
		CRS	1.0%	<u>0.9331</u>	0.8444	0.8777	<u>0.9715</u>	0.9241	0.9358	<u>9.5913</u>	14.1650	11.6503
(1000, 3, 3)	Rad	VRS	8.6%	<u>0.8971</u>	0.7272	0.6575	<u>0.9550</u>	0.8589	0.8178	<u>0.0361</u>	0.0519	0.0539
		CRS	3.8%	<u>0.9104</u>	0.7551	0.7394	<u>0.9611</u>	0.8721	0.8636	<u>0.0378</u>	0.0563	0.0539
	Dir	VRS	8.6%	<u>0.9031</u>	0.7356	0.6530	<u>0.9547</u>	0.8539	0.7986	<u>8.6571</u>	12.7943	13.4832
		CRS	3.8%	<u>0.9228</u>	0.8068	0.7693	<u>0.9677</u>	0.9034	0.8721	<u>9.1480</u>	14.1488	13.1869
(1000, 5, 5)	Rad	VRS	31.6%	<u>0.8467</u>	0.5957	0.4888	<u>0.9202</u>	0.7325	0.6842	<u>0.0314</u>	0.0484	0.0485
		CRS	24.4%	<u>0.8473</u>	0.5459	0.5032	<u>0.9255</u>	0.7137	0.6927	<u>0.0281</u>	0.0451	0.0432
	Dir	VRS	34.5%	<u>0.8199</u>	0.5418	0.4307	<u>0.8873</u>	0.6802	0.6181	<u>5.4077</u>	8.3410	8.3921
		CRS	23.5%	<u>0.8464</u>	0.5779	0.5406	<u>0.9229</u>	0.7325	0.7153	<u>6.4614</u>	10.0452	9.5070
(3000, 2, 2)	Rad	VRS	1.0%	<u>0.9374</u>	0.8659	0.7944	<u>0.9706</u>	0.9373	0.9144	<u>0.0320</u>	0.0426	0.0507
		CRS	0.4%	<u>0.9509</u>	0.8930	0.8266	<u>0.9751</u>	0.9484	0.9189	<u>0.0313</u>	0.0447	0.0544
	Dir	VRS	1.1%	<u>0.9606</u>	0.8935	0.8226	<u>0.9817</u>	0.9490	0.9077	<u>8.4381</u>	12.9187	14.5997
		CRS	0.3%	<u>0.9579</u>	0.9152	0.9183	<u>0.9796</u>	0.9612	0.9568	<u>12.2077</u>	16.3358	14.1900
(3000, 3, 3)	Rad	VRS	3.8%	<u>0.9319</u>	0.8174	0.7131	<u>0.9703</u>	0.9152	0.8595	<u>0.0288</u>	0.0417	0.0475
		CRS	1.7%	<u>0.9369</u>	0.8144	0.7659	<u>0.9704</u>	0.9059	0.8840	<u>0.0326</u>	0.0487	0.0511
	Dir	VRS	3.6%	<u>0.9442</u>	0.8444	0.7278	<u>0.9740</u>	0.9246	0.8522	<u>7.7219</u>	11.7569	13.3009
		CRS	1.6%	<u>0.9506</u>	0.8655	0.8193	<u>0.9774</u>	0.9346	0.9041	<u>7.9798</u>	11.8693	11.8587
(3000, 5, 5)	Rad	VRS	20.2%	<u>0.8889</u>	0.6508	0.5155	<u>0.9523</u>	0.8015	0.7047	<u>0.0252</u>	0.0401	0.0433
		CRS	12.2%	<u>0.9099</u>	0.6682	0.5973	<u>0.9625</u>	0.8165	0.7669	<u>0.0274</u>	0.0460	0.0455
	Dir	VRS	19.8%	<u>0.9084</u>	0.6823	0.5166	<u>0.9560</u>	0.8120	0.6974	<u>5.3507</u>	9.0719	9.9465
		CRS	13.2%	<u>0.9108</u>	0.6883	0.5943	<u>0.9604</u>	0.8230	0.7543	<u>5.2343</u>	8.8918	8.9578
(5000, 2, 2)	Rad	VRS	0.7%	<u>0.9421</u>	0.8917	0.7697	<u>0.9694</u>	0.9525	0.9014	<u>0.0362</u>	0.0462	0.0605
		CRS	0.2%	<u>0.9453</u>	0.8988	0.8301	<u>0.9705</u>	0.9495	0.9195	<u>0.0291</u>	0.0414	0.0536
	Dir	VRS	0.6%	<u>0.9627</u>	0.9277	0.8323	<u>0.9816</u>	0.9674	0.9124	<u>9.3003</u>	12.9303	15.7681
		CRS	0.2%	<u>0.9628</u>	0.9217	0.9112	<u>0.9814</u>	0.9636	0.9533	<u>9.7827</u>	13.7791	12.8037
(5000, 3, 3)	Rad	VRS	2.7%	<u>0.9380</u>	0.8339	0.7132	<u>0.9716</u>	0.9237	0.8618	<u>0.0306</u>	0.0441	0.0516
		CRS	1.1%	<u>0.9517</u>	0.8481	0.7826	<u>0.9783</u>	0.9272	0.8939	<u>0.0278</u>	0.0434	0.0485
	Dir	VRS	2.9%	<u>0.9504</u>	0.8184	0.7208	<u>0.9767</u>	0.9057	0.8490	<u>6.4999</u>	10.0265	11.4096
		CRS	1.2%	<u>0.9471</u>	0.8520	0.8067	<u>0.9741</u>	0.9256	0.8974	<u>8.1962</u>	12.0472	12.3477
(5000, 5, 5)	Rad	VRS	15.6%	<u>0.9097</u>	0.6884	0.5412	<u>0.9637</u>	0.8318	0.7279	<u>0.0243</u>	0.0404	0.0444
		CRS	8.5%	<u>0.9314</u>	0.7200	0.6260	<u>0.9719</u>	0.8480	0.7873	<u>0.0242</u>	0.0436	0.0448
	Dir	VRS	15.6%	<u>0.9235</u>	0.7117	0.5438	<u>0.9662</u>	0.8380	0.7195	<u>5.2906</u>	9.1541	10.1447
		CRS	8.7%	<u>0.9325</u>	0.7413	0.6528	<u>0.9708</u>	0.8600	0.7988	<u>5.3279</u>	9.1580	9.3668

나은 성능을 보이는 것도 주목할만하다. Density에 따른 차이는 크게 발견되지는 않으나 density가 커질수록 학습성능이 다소 떨어지는 현상은 나타났다.

기계학습 모형 간의 성능차이도 확인이 되는데, 거의 대부분의 경우에서 SVM 또는 SVR이 다른 방법들보다 비교적 큰 우위를 보이고 있다. 즉, DEA 모형의 효율적 경계선을 식별하는 데는 SVM 또는 SVR 모형이 상대적으로 더 효과적임을 나타내고 있다. 이러한 경향은 데이터의 크기가 커질수록 더 두드러진다. 랜덤포레스트 모형의 경우 SVM 또는 SVR 모형에 비해 낮은 성능을 보이지만, 전통적인 로지스틱회귀분석이나 선형회귀분석에 비해서는 상대적인 우수성을 보였다.

한편, 훈련데이터의 분할비율을 50%에서 70%로 증가시키는 경우 예측성능이 어떤 영향을 받는지 확인하였다. 지면 제약 관계상 여기서는 DMU 개수가 $n=3000$, 투입-산출요소 개수가 $m=s=3$, VRS 방사형 DEA 모형을 대상으로 한 실험 결과만을 <Table 3>과 <Table 4>에 제시하며, 그 외의 경우는 이와 유사한 결과를 얻었다는 점을 밝힌다.

상기 실험 결과를 보면 전체 모집단으로부터 추출되는 훈련데이터의 비율을 높이면 학습모형의 예측성능이 대체로 높아지는 경향을 확인할 수 있다. 다만, 훈련데이터 비율의 증가에 따른 예측성능의 개선폭은 상대적으로 둔감하게 반응하는 것을 관찰할 수 있고, 일부 경우는 개선이 이루어지지 못하거나 오히려 나빠지는 경우도 있었다. 또한, 훈련데이터의 양이 늘수록 계산량이 함께 증가하므로, 그 효과와의 트레이드오프를 검토하여 적정

한 비율에 대한 의사결정이 이루어져야 한다.

4.4 풀이 시간 비교

제안하는 방법은 빅데이터 환경 하에서 풀어야 하는 DEA 모형의 크기가 상당한 규모로 커짐에 따라 계산 속도가 급격히 느려져 실시간 평가를 요하는 상황에서 현실적 어려움이 발생할 수 있다는 점을 인식하고 이를 극복하기 위해 DEA와 기계학습 모형을 결합하여 효율성 추정치의 계산량을 줄이려는 목적을 가진다. 따라서 이러한 목적을 달성했는지 여부를 실제 계산시간을 통해 확인할 필요가 있다. 여기서는 주어진 데이터 전체를 대상으로 DEA 모형을 푸는 통상적인 방법의 계산시간과 제안하는 방법을 통해 효율성을 추정하는 방법의 계산시간을 비교한다. 각 데이터 규모별로 두 가지 방식의 계산시간을 초 단위로 측정하고 이들 간의 비율을 계산한 결과는 <Table 5>와 같다. 단, 이 실험에서는 VRS 방사형 DEA 모형과 함께 기계학습 모형으로 본 연구에서 가장 우수한 예측성능을 보인 SVR 모형을 사용하였고, 훈련데이터의 비율은 50%이다. 실험을 위해 사용한 PC의 사양은 CPU AMD Ryzen 7 2700X 3.70Ghz, RAM 16GB, OS Windows 10 Pro이다.

실험 결과를 통해 볼 수 있듯이, 데이터 규모가 커질수록 제안하는 방법을 통한 계산시간이 전체 데이터를 대상으로 DEA 모형을 푸는 통상적인 방법에 비해 크게 줄어드는 것을 확인할 수 있다. 특히, 데이터의 크기가 커질

Table 3. Classification Performance when $\rho = 0.7$

Size (n, m, s)	DEA Model		Density	Accuracy			F1			Balanced accuracy		
				SVM	RF	LoR	SVM	RF	LoR	SVM	RF	LoR
(3000, 3, 3)	Rad	VRS	3.8%	0.9750	0.9698	0.9628	0.6779	0.5695	0.4096	0.8649	0.7754	0.6759
		CRS	1.7%	0.9875	0.9855	0.9829	0.6465	0.4820	0.4313	0.8548	0.7173	0.7056
	Dir	VRS	3.6%	0.9726	0.9661	0.9557	0.6883	0.5807	0.3420	0.8615	0.7742	0.6333
		CRS	1.6%	0.9884	0.9854	0.9834	0.6386	0.4503	0.4213	0.8131	0.6877	0.6805

Table 4. Regression Performance when $\rho = 0.7$

Size (n, m, s)	DEA Model		Density	R^2			Spearman correlation			RMSE		
				SVR	RF	LR	SVR	RF	LR	SVR	RF	LR
(3000, 3, 3)	Rad	VRS	3.8%	0.9506	0.8475	0.7164	0.9799	0.9319	0.8629	0.0231	0.0382	0.0482
		CRS	1.7%	0.9639	0.8547	0.7823	0.9842	0.9306	0.8940	0.0203	0.0387	0.0448
	Dir	VRS	3.6%	0.9592	0.8310	0.7013	0.9818	0.9132	0.8356	4.2350	7.9380	9.8739
		CRS	1.6%	0.9650	0.8732	0.8244	0.9844	0.9372	0.9071	5.3343	9.9798	10.8238

Table 5. Computational Time Improvement of the Proposed Approach

Size (n, m, s)	Computation time (seconds)		Ratio (%) ② / ①
	① DEA with all data	② DEA with training data + SVR	
(1000, 2, 2)	0.4113	0.1414 (0.1221 + 0.0193)	34.4
(1000, 3, 3)	0.6628	0.2041 (0.1759 + 0.0282)	30.8
(1000, 5, 5)	1.1676	0.3268 (0.2884 + 0.0384)	28.0
(3000, 2, 2)	3.7658	1.0414 (0.9719 + 0.0695)	27.7
(3000, 3, 3)	6.2002	1.6117 (1.4778 + 0.1339)	26.0
(3000, 5, 5)	12.0139	3.1085 (2.8258 + 0.2827)	25.9
(5000, 2, 2)	10.0035	2.6040 (2.4636 + 0.1404)	26.0
(5000, 3, 3)	17.3081	4.4543 (4.1584 + 0.2959)	25.7
(5000, 5, 5)	47.8801	12.3843 (11.1475 + 1.2368)	25.9

수록 이러한 효과는 점차 심화되고, DMU의 개수가 5,000 개에 이르는 경우에는 계산시간이 거의 1/4 까지 떨어지는 것을 볼 수 있다.

4.5 데이터 증강의 효과

본 연구에서 제안한 데이터 증강 방법에 대한 효과를 살펴보기 위해 데이터 증가 여부 및 수준에 따른 기계학습 모형의 예측성능 차이를 비교하였다. 지면 제약 관계상 여기서는 DMU 개수가 $n=3000$, 투입-산출요소 개수가 $m=s=3$, VRS 방식형 DEA 모형을 대상으로 한 실험 결과만을 제시하며, 그 외의 경우는 이와 유사한 결과를 얻었다는 점을 밝힌다.

데이터 증강을 실시한 실험결과와는 앞서 이미 제시하였으므로 여기서는 데이터 증강을 실시하지 않은 실험결과를 제시하며, 분류모형과 회귀모형으로 나누어 각각 세 가지 학습모형의 예측성능 평가척도를 계산한 결과는 <Table 6>과 <Table 7>과 같다.

상기 실험 결과에서 확인할 수 있듯이, 데이터 증강은 학습모형의 예측 성능을 개선시키는데 일부 도움을 주는 것을 알 수 있다. 특히, 주목할 만한 부분은 분류모형에서의 F1 점수와 균형정확도에서 큰 개선이 이루어졌다는 점이다. 본 연구에서 제안한 데이터 증강 방법이 효율적-비효율적 데이터 클래스 간의 심한 불균형을 개선하고 효율적 경계선을 보다 뚜렷하게 나타내는 훈련데이터를 구성하는데 목적이 있는데, 상기 결과는 이러한 목적이 적절하게 달성되었음을 보여준다

Table 6. Classification Performance without Data Augmentation

Size (n, m, s)	DEA Model		Den-sity	Accuracy			F1			Balanced accuracy		
				SVM	RF	LoR	SVM	RF	LoR	SVM	RF	LoR
(3000, 3, 3)	Rad	VRS	3.8%	0.9726	0.9645	0.9627	0.5517	0.5228	0.3911	0.7219	0.7373	0.6509
		CRS	1.7%	0.9845	0.9843	0.9841	0.1679	0.3775	0.4072	0.5313	0.6548	0.6707
	Dir	VRS	3.6%	0.9715	0.9648	0.9618	0.5457	0.5462	0.3916	0.7231	0.7557	0.6513
		CRS	1.6%	0.9834	0.9829	0.9838	0.1657	0.3773	0.4036	0.5325	0.6514	0.6656

Table 7. Regression Performance without Data Augmentation

Size (n, m, s)	DEA Model		Den-sity	R^2			Spearman correlation			RMSE		
				SVR	RF	LR	SVR	RF	LR	SVR	RF	LR
(3000, 3, 3)	Rad	VRS	3.8%	0.9319	0.8127	0.7076	0.9714	0.9088	0.8554	0.0317	0.0448	0.0527
		CRS	1.7%	0.9302	0.7947	0.7676	0.9704	0.8916	0.8870	0.0284	0.0425	0.0437
	Dir	VRS	3.6%	0.9422	0.7942	0.7123	0.9729	0.8856	0.8431	6.1336	9.5218	11.1023
		CRS	1.6%	0.9457	0.8380	0.8012	0.9762	0.9170	0.8943	6.8388	10.4904	10.9691

5. 결 론

표준적인 DEA 모형은 기본적으로 선형계획모형을 풀어 효율성을 분석할 수 있고, 선형계획모형은 상대적으로 쉬운 문제이며 이를 빠른 시간에 풀어낼 수 있는 고급 상용소프트웨어들이 많이 개발되어 있다. 또한, 대규모 DEA 모형을 효율적으로 신속하게 풀기 위한 다양한 형태의 수리계획 기반 방법들이 그간 폭넓게 개발되어 왔다. 그러나 최근 들어 빅데이터의 시대가 도래하면서 기존에 상상하기 어려웠던 초대형 데이터가 수집되고 있고, 이러한 데이터를 거의 실시간으로 분석해내야 하는 응용들이 나타나고 있어 계산의 신속성 문제가 다시 한번 중요해지고 있다. DEA 또한 최근 들어 기존의 전통적인 생산운영 프로세스의 효율성을 측정하는 도구라는 인식을 넘어 데이터 기반 애널리틱스(Data Enabled Analytics) 방법론으로 자리매김하려는 시도들이 대두되고 있다. 따라서 빅데이터 환경 하에서 빠르게 DEA 효율성 분석을 해낼 수 있는 방법론의 필요성이 높아지고 있다.

본 연구에서는 빅데이터 DEA 모형의 효율적 계산을 위한 방법론으로 기계학습에 기반한 방법을 제안하였다. 대규모 모집단 DMU 전체를 대상으로 DEA 모형을 푸는 방법 대신, 일부 훈련데이터를 추출하여 DEA 효율적 경계선과 그로부터의 거리를 학습하는 학습모형을 훈련시킨 후 이를 활용하여 훈련데이터 이외 나머지 DMU에 대한 효율성 점수를 추정하거나 새롭게 스트리밍되어 유입되는 DMU에 대한 효율성 점수를 예측하는 방법을 제안하였다. 이와 더불어, 효율적 DMU와 비효율적 DMU 간의 데이터 클래스 불균형 문제를 해소하고 효율적 경계선을 보다 분명하게 드러내는 훈련데이터를 구성할 수 있도록 비효율적 DMU를 완전 또는 불완전하게 투영하는 방식의 데이터 증강법을 제시하였다.

새롭게 제안하는 방법에 대한 유효성을 실증하기 위해 다양한 DEA 모형, 여러 가지 기계학습 방법론, 다양한 데이터 규모 등을 조합하여 일정한 규칙에 따라 랜덤하게 생성된 데이터를 대상으로 효율성 점수 예측 성능을 분석하였다. 실증 실험 결과 제안한 방법론의 예측 성능이 상당한 수준으로 나타났으며, 특히 SVM 또는 SVR을 학습모형으로 활용하였을 때 효율적 경계선 식별 및 효율성 점수 예측 성능이 가장 우수한 것으로 나타났다. 모집단 DMU 전체를 대상으로 DEA 모형을 직접적으로 푸는 방식에 비해 제안하는 접근법이 가지는 계산 시간 측면의 우수성도 함께 보였다.

REFERENCES

- [1] Barr, R. S., & Durchholz, M. L. (1997). Parallel and hierarchical decomposition approaches for solving large-scale data envelopment analysis models. *Annals of Operations Research*, 73, 339-372.
- [2] Charles, V., Aparicio, J., & Zhu, J. (2019). The curse of dimensionality of decision-making units: A simple approach to increase the discriminatory power of data envelopment analysis. *European Journal of Operational Research*, 279(3), 929-940.
- [3] Charnes, A., Cooper, W. W., & Rhodes, E. (1978). Measuring the efficiency of decision making units. *European Journal of Operational Research*, 2(6), 429-444.
- [4] Chen, W. C., & Cho, W. J. (2009). A procedure for large-scale DEA computations. *Computers & Operations Research*, 36(6), 1813-1824.
- [5] Chen, W. C., & Lai, S. Y. (2017). Determining radial efficiency with a large data set by solving small-size linear programs. *Annals of Operations Research*, 250(1), 147-166.
- [6] Dulá, J. H. (2008). A computational study of DEA with massive data sets. *Computers & Operations Research*, 35(4), 1191-1203.
- [7] Dulá, J. H. (2011). An algorithm for data envelopment analysis. *INFORMS Journal on Computing*, 23(2), 284-296.
- [8] Dulá, J. H. & Helgason, R. V. (1996). A new procedure for identifying the frame of the convex hull of a finite collection of points in multidimensional space. *European Journal of Operational Research*, 92(2), 352-367.
- [9] Dulá, J. H. & López, F. J. (2013). DEA with streaming data. *Omega*, 41(1), 41-47.
- [10] Dulá, J. H. & Thrall, R. M. (2001). A computational framework for accelerating DEA. *Journal of Productivity Analysis*, 16(1), 63-78.
- [11] Farrell, M. J. (1957). The measurement of productive efficiency. *Journal of the Royal Statistical Society: Series A (General)*, 120(3), 253-281.
- [12] Khezrimotlagh, D., Zhu, J., Cook, W. D., & Toloo, M. (2019). Data envelopment analysis and big data. *European Journal of Operational Research*, 274(3), 1047-1054.
- [13] Koopmans, T. C. (1951). *Activity Analysis of Production and Allocation*, Wiley, New York.
- [14] Lim, S., Oh, K. W., & Zhu, J. (2014). Use of DEA cross-

- efficiency evaluation in portfolio selection: An application to Korean stock market. *European Journal of Operational Research*, 236(1), 361-368.
- [15] Pareto, V. (1927). *Manuel d'economie politique*, deuxieme edition, Appendix, pp. 617 ff., Alfred Bonnet, ed., Marcel Giard, Paris.
- [16] Shi, Y., Zhu, J., & Charles, V. (2020). Data science and productivity: A bibliometric review of data science applications and approaches in productivity evaluations. *Journal of the Operational Research Society*, 72(5), 975-988.
- [17] Stewart, T. J. (1996). Relationships between data envelopment analysis and multicriteria decision analysis. *Journal of the operational research society*, 47(5), 654-665.
- [18] Sueyoshi, T. & Chang, Y. L. (1989). Efficient algorithm for additive and multiplicative models in data envelopment analysis. *Operations Research Letters*, 8(4), 205-213.
- [19] Zhu, J. (2020). DEA under big data: Data enabled analytics and network data envelopment analysis. *Annals of Operations Research*, 1-23.
- [20] Zhu, Q., Wu, J., & Song, M. (2018). Efficiency evaluation based on data envelopment analysis in the big data context. *Computers & Operations Research*, 98, 291-300.



응위엔투이즈영

노동사회대학교(베트남) 경영학과 학사
 현재: 동국대학교 일반대학원 경영학과 석사과정
 관심분야: 효율성 평가 및 벤치마킹, 데이터과학



임 성 목

서울대학교 산업공학과 학사
 서울대학교 산업공학과 석사
 서울대학교 산업공학과 박사
 현재: 동국대학교 경영대학 경영학과 교수
 관심분야: 효율성 평가 및 벤치마킹, 비즈니스애널리틱스, 정보화

Industry 4.0 기술과 디지털 SCM 실행방식의 연계: 중국 시장을 중심으로

김영길* · 박정수**†

*신한대학교 글로벌통상경영학과, **중앙대학교 다빈치 교양대학

Digital Supply Chain Management and Performance in Chinese Companies: The Role of Artificial Intelligence Big Data Usage

Yeonggil Kim* · Jeong Soo Park**†

*Department of Global Trade and Management, Shinhan University

**Da Vinci College of General Education, Chung-Ang University

This study aims to check whether Digital Supply Chain Management practices has positive effect on company's performance by using sample consisting of companies in Chinese industry. Furthermore, the authors check if companies' utilizing either AI related technologies or Big Data analytics have positive moderating or accelerating effect on the relationship between DSCM and the performance. To achieve these goals, the researchers conducted surveys and empirical researches using validity check, reliability analysis, regression model and moderate regression model. Results of further research showed that DSCM practices affects the performance positively in Chinese sample companies. As the further research results, utilizing AI related technologies has positive moderating effect on relationship between DSCM practices and the performance, while utilizing Big Data analytics does in limited range. From these research results a managerial implication that implementing DSCM practices and AI related technologies at the same time gives Chinese companies more improvement in performances due to the synergetic effects between them.

Keyword : Digital Supply Chain Management, Artificial Intelligence, Big Data Analytics, Performance, Moderate Regression

† Corresponding Author : 84, Heukseok-ro, Dongjak-gu, Seoul, 06974, Korea, Tel: +82-2-820-6798, E-mail: poshboy@cau.ac.kr

Received: 19 November 2021, Revised: 8 December 2021, Accepted: 13 December 2021

1. 서 론

고객 서비스 개선에 의한 만족과 전체 비용 감소를 목표로 수직적으로 연결된 다수 업체들의 협력과 협업(collaboration)을 통하여 가시성과 통제가능성을 높여 결과적으로 동시화(synchronization)의 달성을 목표로 하는 공급사슬관리(SCM)은 다양한 방향에서 접근이 이루어졌으며, 그 연구 방향은 공급사슬의 구성요소, 공급사슬 관리 프로세스, 공급사슬 내 물적·수요정보·재무적 등의 흐름의 관리, 공급사슬 네트워크 구조의 네 가지로 요약된다(Garay-Rondero et al., 2020). 최근 이러한 공급사슬관리에 새로운 흐름으로서 등장한 것이 디지털 공급사슬관리(DSCM)로서 이는 위에서 언급된 공급사슬관리의 목표 달성을 위하여 정보 기술은 물론 다양한 새로운 기술 특히 4차 산업혁명 혹은 Industry 4.0의 구성요소 혹은 기반 기술로 통칭되는 혁신적 기술들을 수용 및 응용하는 것을 특징으로 한다. 최근 수많은 관심이 집중되고 있는 4차 산업혁명 혹은 Industry 4.0은 다양한 혁신적 기술 특히 ICT가 다양한 산업들과 결합하여 새로운 형태의 제품, 서비스, 사업들을 창출하여 개인과 사회 전반에 심도깊은 변화가 발생하고 있음을 뜻하는데, 과거 기술과는 달리 실제와 가상을 통합하고 자동적으로 그리고 지능을 활용하여 제어하는 기술들을 포괄한다는 특성을 지닌다. 이러한 4차 산업혁명에 활용되는 기술들은 미래 신성장동력으로 불리기도 하며, 그 종류가 무궁무진하지만 대표적인 것으로 인공지능(Artificial Intelligence), 사물인터넷(IoT), 빅데이터 분석(Big Data Analytics), 블록체인(Blockchain), 클라우드 컴퓨팅(Cloud Computing), 드론(Dron or UAV), 증강 및 가상 현실(AR/VR), 로봇(Robot)기술 등은 공통적으로 포함되며 여기에 스마트 공장, 스마트 시티 등을 추가하여 최종적으로는 가상물리시스템(Cyber Physical System)을 구축하여 인간의 업무와 문화를 혁신적으로 변화시킨다는 목표를 지니고 있다.

이러한 디지털 공급사슬관리의 개념 및 전략과 4차 산업혁명 관련 기술에 대한 관심은 폭발적으로 증가하고 있으나, 그러한 관심에 비하여 국내의 연구는 그리 많지 않으며 특히 양자를 연계한 연구 및 그들과 기업의 성과와의 관련성을 검토하는 연구는 찾아보기 쉽지 않은 것이 현재의 상황이다.

본 연구는 이러한 상황에서 벗어나 미래지향적이고 최신의 추세를 반영하는 연구로서 디지털 공급사슬관리의 실행과 조직 특히 기업의 성과와의 관련성 여부를

검증하고 그 다음 단계로 현재 많은 관심이 주어지고 있는 4차 산업혁명 혹은 Industry 4.0의 다양한 기반 기술 중 인공지능(AI)과 빅데이터 분석의 적극적 실행이 디지털 공급사슬관리 실행과 성과 사이의 관계에 대하여 조절적 역할 즉 그 관계를 강화시키는 역할을 수행하는지를 실증적으로 검토하고 검증하는 것을 연구 목적으로 하였다.

보다 구체적으로 본 연구는 첫째로 거대한 시장이라고 간주되는 중국 내에 위치한 다수 업체들을 연구 표본으로 하여 디지털 공급사슬관리 실행이 그들의 성과에 정(+)의 영향을 미치고 있는가를 서베이를 통하여 검증하는 것을 달성해야 할 연구목표로 하였다. 다음으로 디지털 공급사슬관리 실행과 성과 간 관계에 대해 AI의 실행과 빅데이터 분석의 실행이 각기 정(+)의 조절적 효과를 나타내는가를 검증하고자 하였다.

본 연구는 진행 순서는 다음과 같다. 두 번째 장에서 본 연구의 연구 과제와 관련된 기존의 문헌과 선행 연구들을 검토한 후, 세 번째 장에서는 본 연구에서 검증하는 것을 목적으로 하는 연구 모형과 연구 가설을 설정 및 제시하기로 한다. 네 번째 장에서는 설정된 연구 모형과 연구 가설들을 서베이 및 실증적 분석 모형과 방법들을 통하여 검증하며, 최종적으로 다섯 번째 장에서는 본 연구의 연구 결과들이 요약 및 제시되고 그 의미와 시사점 그리고 연구의 공헌들이 제시된다.

2. 기존 관련 문헌 연구

2.1 디지털 공급사슬관리

먼저 디지털 공급사슬관리(Digital Supply Chain Management: DSCM)의 정의를 살펴보기로 한다. 기존의 관련된 연구들을 정리하여 연대기적으로 분류한 Buyukozkan and Gocer(2018)에 따르면, DSCM이란 제품의 생산 및 서비스의 제공 과정에서 보다 가치를 향상시키고 보다 편리하게 구입할 수 있도록 그리고 보다 유연성을 높이는 것을 가능하게 하는 공급사슬의 새로운 접근이며, SCM의 구성원 사이의 다양한 활동을 지원하고 동시화(synchronize)하기 위하여 대량 데이터 처리, 협력, 디지털 하드웨어, 소프트웨어, 네트워크를 통한 커뮤니케이션 능력에 기반을 둔 지능적인 시스템으로 정의될 수 있다. 동 연구는 다양한 DSCM의 기존 문헌들에 기초하여 그 특징으로서 신속한 속도, 높은 유연성, 범세계적 차원의 연결, 실시간 기초 재고관리, 지능적인 관리, 투명

성 제고, 측정가능성 향상, 혁신성의 제고, 선제적 문제 해결, 환경친화성이 제고를 들었다.

Buyukozkan and Gocer(2018) 이전의 DSCM의 정의의 하나로, Bhargawa et al.(2013)은 그것을 범세계적 차원으로 분산되어 있는 조직들 사이의 활동을 지원하고 공급사슬 내 행위자들 사이의 활동을 이끌어 가는 시스템으로 정의하고, 이것에는 소프트웨어와 하드웨어는 물론 커뮤니케이션 네트워크 모두가 포함되며, 그러한 활동들에는 구매, 제조, 보관, 운송, 판매 등이 포함된다고 설명하였다. 최근 널리 인용 및 언급되는 DSCM의 정의의 하나로, Kinnet(2015)는 그것을 새로운 형태의 수익 및 가치를 창출하기 위한 기술과 분석 기법(analytics)을 활용하는 새로운 기법들을 보다 적극적으로 적용하는(leverage) 지능적인 그리고 가치를 창출하는 네트워크라고 정의하였다.

디지털 공급사슬관리(DSCM)의 선행적 개념을 활용한 연구로서 정보 기술(IT)을 공급사슬관리에 적용한 선구적 연구 중 하나로 볼 수 있는 Rai et al.(2006)는 IT의 통합적 하부구조가 기업의 공급사슬 프로세스 통합 역량에 영향을 미치는지 그리고 그 역량이 기업 성과에 긍정적인 영향을 미치는가에 대한 연구를 실행하였다. 이 연구는 SCM의 효과적 실행을 위한 IT의 통합적 하부구조에 데이터 일관성 및 다기능적 응용능력을 포함시키고 공급사슬 프로세스 역량에는 정보흐름의 통합, 물적 흐름의 통합, 재무적 흐름의 통합을 포함시켰으며 기업 성과에는 운영 측면의 탁월성 개선, 고객 관계 개선 수익 증가를 포함시키고 있다. 이후의 연구로서 Li et al.(2009)는 IT의 실행이 공급사슬 통합과 공급사슬성과에 각각 영향을 미치는지 그리고 공급사슬 통합이 다시 공급사슬의 성과에 긍정적 영향을 미치는가를 구조 방정식 모형에 의하여 실증적으로 검토하였다. 이들은 IT의 실행에 전자문서교환(EDI)의 활용, 바코드의 활용, 개방적 표준과 특수식별 코드의 적용, 협력업체에 대한 지원과 의사결정 시스템의 활용을 포함시켰고 공급사슬 통합에는 물류중심 설계(Design for logistics)에 기반을 둔 물류시스템 자원의 최적화 전략의 실행, 시장 추세의 이해와 정확한 수요 예측, SCM 계획의 정확성 및 적응성, 재고 물품의 통제 및 추적에서의 정확성과 가시성 제고, 프로세스의 표준화와 가시성 제고를 포함시켰다. 한편, 공급사슬 성과에는 적시 생산성과, 재고 회전율과 현금 회전율, 고객 조달 기간과 화물 운송 효율성, 배송 측면의 성과와 품질 향상, 공급사슬 재고의 가시성 증가와 기회비용 감소, 총물류비용 감소를 포함시켰다. 정보기술과 SCM을 연계시킨 또다른 연구인 Wong et al.(2011)에 따르면 정보기술에

의한 공급사슬의 통합은 기업의 다양한 운영 활동들을 보다 적절히 조정하는 데 도움을 제공함으로써 공급사슬 관리의 활동들이 보다 원활하게 수행될 수 있게 하는 역할을 수행한다.

Xue et al.(2013)은 시스템의 모듈화(modularity)가 디지털 공급사슬의 인식된 위험에 미치는 영향과 후자가 다시 공급사슬 디지털화(digitization)에 미치는 영향을 검토함으로써 IT 거버넌스가 조절 효과를 갖는지를 검증하는 연구를 수행하였다. 이들은 동 연구에서 DSCM과 유사한 개념으로서 ‘공급사슬 디지털화’를 제시하였고 그것에 디지털 공급사슬 시스템을 통하여 거래하는 공급업체의 수 및 거래량 그리고 디지털 공급사슬 시스템을 통하여 거래하는 소비자 즉 고객의 수와 거래량을 포함시켰다.

Liu et al.(2013)은 기업의 IT 관련 능력이 성과에 미치는 영향력에 대하여 흡수적 역량과 공급사슬 유연성(agility)이 조절적 효과를 갖는지를 검증하는 연구를 수행하였다. 이 연구에서는 공급사슬관리의 실행 및 역량과 비슷한 개념으로서 SC 유연성을 설정하였는데, 그것에는 ‘가시성’, ‘공동 계획’, ‘프로세스 통합’, ‘공유된 가치’의 네 가지 요인이 포함된다. 이 연구의 후속 연구로 간주될 수 있는 Liu et al.(2015)은 공급사슬 상에서 기업의 IT 운영 역량과 IT 변환 역량이 ‘인터넷에 기반을 둔 공급 통합’과 ‘인터넷에 기반을 둔 수요 통합’을 거쳐 기업 성과에 미치는 영향을 경로분석을 통하여 검증하는 연구를 제안하였다. 이 연구에서는 인터넷에 기반을 둔 공급 통합에는 협력 업체와의 재고 관련 조정 계획 및 정보의 공유, 주문 관련 정보의 공유, 구매 관련 주문의 공유가 포함되었고 인터넷에 기반을 둔 수요 통합에 협력 업체와의 공동 수요 예측, 새로운 제품/서비스 도입 계획의 공유, 서비스 관련 지원 활동의 공유가 포함되었다.

최근의 연구로서 Acimovic and Stajic(2019)은 디지털 공급사슬관리(DSCM)를 공급사슬 네트워크 전체의 성과 개선을 달성하기 위하여 혁신적이고 지능적으로 작동하는 기술적 솔루션에 기반을 두고 공급사슬 내에서 작동되는 비즈니스 모델로 정의하였다. 이들에 의하면 DSCM 실행에 의하여 공급사슬 구성원들은 투명성, 의사소통, 협업, 유연성, 대응성 등의 측면에서의 향상을 기대할 수 있게 된다.

Garay-Rondero et al.(2020)은 DSCM에 대한 방대한 양의 과거 연구들에서 관련된 4차 산업혁명 기술들을 정리하고 분석한 후 디지털 공급사슬관리의 모형을 제시하였

다. 이들은 공급사슬관리의 기존 연구 흐름을 공급사슬 관리의 구성요소에 대한 연구, 공급사슬관리 프로세스에 대한 연구, 공급사슬 내의 다양한 흐름(물적 흐름, 정보 흐름, 재무적 흐름 등)에 대한 연구 그리고 공급사슬 네트워크 구조에 대한 연구의 네 가지 범주로 분류하였다. 이러한 분류에 기초하여 이들은 가상 가치 사슬의 기반 위에 디지털 공급사슬의 구축과 실행이 요구된다고 제안하였으며, 디지털 공급사슬관리의 하위 범주로서 디지털 네트워크 구조, 디지털 공급사슬 흐름, Industry 4.0의 다양한 기술들, 디지털 공급사슬관리 프로세스, 디지털 공급사슬 구성요소를 제시하였다.

2.2 AI와 빅데이터를 포함하는 4차 산업혁명의 구성 요소 및 신기술

4차 산업혁명 그리고 해외에서는 Industry 4.0의 구성 요소로 꼽히는 또한 DSC에서 활용되는 혁신적 기술로 간주되는 최근의 기술적 발전에는 증강현실(Augmented Reality), 가상 현실(Virtual Reality), 빅 데이터(Big Data), 클라우드 컴퓨팅(Cloud Computing), 로봇기술(Robotics), 센서 기술(Sensor Technology), 옴니 채널(Omni Channel), 사물인터넷(Internet of Things), 자율운행차(Self-Driving Vehicles), 드론(Drone or Unmanned Aerial Vehicle), 나노기술(Nanotechnology), 3D 프린팅(3D printing), 등이 포함된다(Buyukozkan and Gocer, 2018).

Barata et al.(2018)은 Industry 4.0과 관련된 공급사슬 관리를 모바일(mobile) SCM으로 정의하고 이와 관련된 기존 연구들을 내용분석 기법으로 분류하여 모바일 SCM에 적용되는 기술들에 무선(wireless), 스마트폰/태블릿, 전자태그(RFID), 글로벌 이동통신 시스템(GSM), 기업자원계획(ERP), 위치추적시스템(GPS), 클라우드 컴퓨팅, 증강현실, 사물인터넷을 포함시키고 있다.

Bar et al.(2018)은 Industry 4.0의 기술들이 중소기업의 공급사슬에 미치는 영향을 검토하는 연구에서 최근의 중요성을 갖는 기술들로서 가상물리시스템(CPS), 스마트 생산 기술, 사물인터넷, 전자태그(RFID), 가상 및 증강 현실, AI를 제시하였다.

이러한 신기술들은 제조 패러다임의 새로운 핵심요인으로 간주될 수 있으며, ‘통합적 제조기술’과 ‘지능적 제조 기술’로 구분될 수 있다. 전자에는 3D 프린팅, 로봇기술에 의한 자동화, 물적 흐름 관리가 후자에는 가상 및 증강현실, 사물인터넷의 산업 응용, 가상 물리 시스템(Cyber Physical System)이 포함된다(Linh et al., 2019). Acimovic

and Stajic(2019)은 4차 산업혁명과 관련된 신기술들을 ‘DSC의 선도적 기술’로 명명하고 블록체인(blockchain) 기술, AI, 증강현실, 빅데이터 분석, 드론 혹은 UAV, 자율주행차, 사물인터넷을 포함시켰다. Garay-Rondero et al.(2020)은 DSCM과 관련된 관련 연구들과 관련된 4차 산업혁명 기술들을 정리하여 DSC 모형을 제시하였는데, 디지털 공급사슬의 하위 범주의 하나에 Industry 4.0의 다양한 기술들을 망라하여 포함시켰다. 그것에는 3D 프린팅, 증강 및 가상 현실, 로봇을 포함하는 공장 자동화 및 스마트 공장, 빅데이터 분석 혹은 데이터 마이닝, 블록체인, 클라우드 컴퓨팅, AI 혹은 기계학습, 가상물리시스템, 사이버보안, 디지털화(digitalization), e-비즈니스, 기업자원계획(ERP), 수직적 혹은 수평적 통합, 인간-기계 상호작용, 정보 통합과 공유, 사물인터넷, 옴니체인(omni-chain) 혹은 옴니채널(omni-channel), 전자태그 혹은 RFID, 스마트 공장 및 물류, 웨어러블(wearable) 기기 활용 등을 포함시키고 있다.

이러한 4차 산업혁명 및 DSCM의 다양한 기술 중 본 연구에서 활용될 첫 번째 기술인 AI는 인간의 지적 능력을 컴퓨터 프로그램으로 실현하여 다양한 문제에서 인간과 유사한 해결능력을 갖추도록 하는 기술을 의미한다. 과거에는 자연언어처리, 전문가시스템, 영상 및 음성 인식, 신경망을 주요 분야로 하였으나, 기술의 발전으로 그 중 대부분이 실현되어 현재에는 기계학습, 인공신경망, 딥러닝(Deep Learning)의 세 가지 분야로 대별된다.

본 연구에서 활용될 4차 산업혁명 혹은 Industry 4.0의 두 번째 기술인 빅데이터 분석은 정보 기술의 발전에 의한 막대한 양의 정보를 관리하고 분석할 뿐 아니라 가치를 추출하는 기법 혹은 분야를 의미한다. 통계학, 기계학습, 인공신경망, 데이터 마이닝의 다양한 기법과 모형을 활용하여 대용량 데이터의 처리와 결과 해석을 통하여 보다 정확한 의사결정과 기업전략의 적절한 설정 및 수정보완을 가능하게 하는 분야로서 중요성을 지닌다(이종호, 2018; Acimovic and Stajic, 2019).

3. 연구모형 및 가설

본 연구는 중국 내 소재 기업들을 연구 표본으로 하여 디지털 공급사슬관리와 기업의 성과가 연구 개념으로서의 적합성을 갖는가를 타당성과 신뢰성 분석을 통하여 검증한 후, 디지털 공급사슬관리가 독립변수로서 기업의 성과라는 종속변수에 대하여 긍정적 영향을 미치는가를 검증하는 것을 첫 번째 연구 주제로 하였다. 다음으로

디지털 공급사슬관리와 기업의 성과 사이의 그러한 관계에 대하여 4차 산업혁명의 다양한 구성요인 중 AI의 실행과 빅데이터 분석의 실행이 각각 조절적 효과를 나타내는가를 검증하는 것을 두 번째 연구 목표로 설정하였다. 이러한 연구의 설계를 그림으로 나타내면 아래의 <Fig. 1>과 같다.

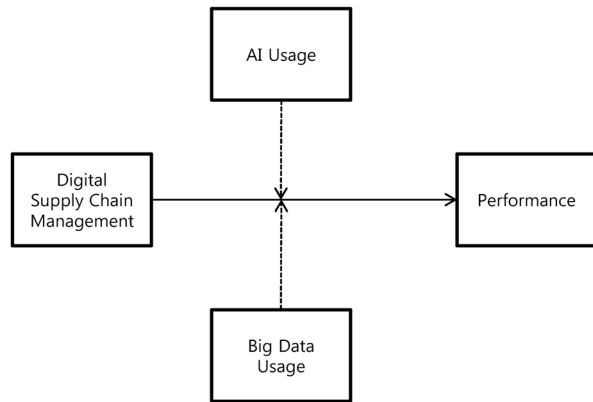


Fig. 1 Conceptual Model of Research

상기한 연구 설계와 연구 모형에 따라서 본 연구에서 적용될 연구 가설들은 아래와 같다.

연구가설 1: 중국 기업들에 있어서 디지털 공급사슬관리 실행은 성과에 정(+)의 영향을 미칠 것이다.

연구가설 2: 중국 기업들의 AI의 적극적 활용은 디지털 공급사슬관리와 성과 간 관계에 대하여 정(+)의 조절 효과를 나타낼 것이다.

연구가설 3: 중국 기업들의 빅 데이터의 적극적 활용은 디지털 공급사슬관리와 성과 간 관계에 대하여 정(+)의 조절 효과를 나타낼 것이다.

디지털 공급사슬관리가 성과에 영향을 미친다는 연구가설 1은 Li et al.(2009)를 참고하였으며 연구가설 2와 3은 기존 문헌 연구의 논문들을 참고하여 저자들에 의하여 독자적으로 설정되었다.

그리고 이러한 AI 및 빅데이터의 활용이 의미를 지니려면 회사의 설립 이후 어느 정도의 연차가 지나고 어느 정도의 규모를 지닌 기업들이어야 한다는 관점에서 연구대상 표본 기업들이 그러한 조건을 만족시키도록 설문대상 기업의 선택에 유의하였다.

4. 실증 분석 결과와 해석

4.1 조사 기업의 인구통계적 특성

본 연구는 중국에 위치한 총 353개 업체들을 대상으로 하여 설문지에 의한 조사가 수행되었으며 실증적 분석이 이루어졌다. 설문은 해당되는 기업의 관리자급 이상을 대상으로 하였고 먼저 응답 관리자들의 근무 연수를 살펴보면 5년 이하가 75명, 6년에서 10년 이하가 90명, 11년 이상에서 15년 이하가 75명, 16년 이상에서 20년 이하가 108명으로 가장 많이 나타났고, 20년 이상이 5명이었다. 한편 표본을 업종 별로 분류하면 건설 업체가 75개, 서비스 업체가 47개, 정보기술 업체가 27개, 의류 업체가 63개, 음료 업체가 14개, 화학 관련 업체가 20개, PC 제조 및 전자 업체가 75개였다.

그리고 회사 설립 이후 연차는 1년 이하에서 5년 이하는 70개, 6년에서 10년 이하는 121개, 11년에서 15년 이하는 72개 업체로 구성하였고 기업 규모 면에서는 중소기업 251개 대기업 102개로 구성하여 특히 AI 및 빅데이터의 활용 경험을 보유하기 위한 연차와 규모를 확보한 기업들이 표본이 될 수 있도록 노력하였다.

Table 1. Sample Classification by Industry

	Item	No.	Ratio
Industry	Construction	75	21.2%
	Service	47	13.3%
	Physical Distribution	32	9.1%
	IT technology	27	7.6%
	Fashion apparel	63	17.8%
	Beverage	14	4.0%
	Chemical	20	5.7%
	PC and electronics	75	21.2%
	Total	353	100%

Table 2. Sample Classification by Employee Number

	Item	No.	Ratio
EmployeeNumber	Less than 20	81	22.9 %
	20 ~ less than 40	83	23.5 %
	40 ~ less than 60	87	24.6 %
	Over 60	102	28.9 %
	Total	353	100%

4.2 변수의 조작적 정의 및 타당성과 신뢰성 검증

본 연구에 사용된 독립변수와 종속변수의 연구개념 (construct or research concept)으로서의 타당성을 검증하기 위해 설문에서 사용된 문항들에 대한 탐색적 요인 분석을 실시하였다. <Table 3>의 첫 다섯 문항은 디지털 공급사슬관리에 대한 문항들이며 그 다음 다섯 문항은 성과에 대한 것들이다. 디지털 공급사슬관리에 대한 문항들로는 4차 산업혁명 기술을 적용한 장기적 관계 설정에의 노력, 그러한 기술을 활용한 신제품 개발에서의 협력, 네트워크 및 시스템을 통한 수요 정보 공유, 정보기술을 적용한 기술적, 재무적 지원의 제공 및 교류와 회의 실시, 4차 산업혁명 기술을 적용한 의사소통, 협력과 그에 기초한 신뢰구축이 포함되었다. 한편 성과에는 최근 2년간의 매출액 증가, 영업이익 증가, 시장 점유율 증가, 투자수익률 증가, 자산이익률 증가가 포함되었다. <Table 3>의 가장 오른쪽 열의 수치와 같이 두 변수 또는 개념 모두 0.7 이상을 나타내어 일반적 기준을 넘는 결과를 나타내었고 두 개념의 연구 변수로서의 신뢰성이 확보되었음이 검증되었다. <Table 3>의 결과와 같이 두 개념 또는 변수 모두 적재(loading)값으로 일반적 기준인 0.6 이상을 넘는 것으로 나타났고 KMO 값은 0.854를 나타내어 두 개념 모두가 연구 개념으로서의 타당성을 갖는 것으로 검증되었다.

다음으로 두 변수 또는 개념들에 대하여 각 설문 문항들의 신뢰도를 측정하기 위해 자주 활용되는 크론바하의 알파 값에 의한 검증이 이루어졌다. <Table 3>의 가장 오른쪽 열의 수치와 같이 두 변수 또는 개념 모두 0.7

이상을 나타내어 일반적 기준을 넘는 결과를 나타내었고 두 개념의 연구 변수로서의 신뢰성이 확보되었음이 검증되었다.

한편, 조절변수로 연구에 적용된 ‘AI 활용’ 및 ‘빅데이터 활용’은 ‘단일’ 문항으로서 측정되었고 실증분석에 활용되었다. 즉 설문지 상에서 ‘해당 기업이 유튜브 등 각종 프로그램의 인공지능 기능과 관련 기술을 적극적으로 활용하는 정도’와 ‘빅데이터 분석 기법을 경영 활동에 적극적으로 활용하는 정도’를 7점 척도로 질문함으로써 측정이 이루어졌다.

4.3 실증적 분석 결과

본 연구의 세 번째 장에서 설정되었던 연구 가설들의 검증을 위하여 회귀분석과 조절적(또는 위계적) 회귀분석에 의한 실증적 검증이 수행되었다. 세 번째 장에서 제시된 바와 같이 첫 번째 연구가설은 독립변수로서 디지털 공급사슬관리 실행이 종속변수로서 성과에 대하여 정(+)의 영향을 미칠 것이라는 가설이었으며 이의 검증을 위하여 단순 회귀분석에 의한 검증이 수행되었다.

첫 번째의 기본적 회귀모형은 독립변수를 디지털 공급사슬관리(DSCM)실행으로 하고 종속변수를 기업의 성과로 하여 단순회귀분석에 의한 검증이 수행되었다. 아래에 있는 <Table 4>의 Step 1의 결과를 토대로 하여 첫 번째 연구가설은 채택되었다. 회귀모형 전체의 유의도를 뜻하는 F 통계량 값은 127.269로서 충분히 큰 값을 나타내어 1% 유의수준에 유의한 결과를 보였다. 한편 회귀분석 모형의 상수항과 디지털 공급사슬관리 변수의 계수는

Table 3. Survey Items' Operational Definition and Confirmatory Factor Analysis Results

Construct	Survey Item/Question	Factor1	Factor 2	Cronbach's α
Digital Supply Chain Management	DSCM1: Trying to maintain long term relationships	.748	.165	.809
	DSCM2: Cooperation in new product development	.735	.321	
	DSCM3: Sharing information of demand and production	.725	.048	
	DSCM4: Mutual support, meeting and conferences	.683	.237	
	DSCM5: Communications, cooperations and trust building	.665	.136	
Performance	PER1: Increase of sales over the last 2 years	.145	.769	.797
	PER2: Increase of operating profit over the last 2 years	.353	.723	
	PER3: Increase of market share over the last 2 years	.095	.677	
	PER4: Increase of return on investment over the last 2 years	.148	.685	
	PER5: Increase of return on asset over the last 2 years	.144	.600	
Measures	Kaiser-Meyer-Olkin measure	0.854		
	Bartlett's Test	1365.077		

1.523과 0.520으로서 유의수준 1%에서 역시 유의한 결과로 볼 수 있는데, 이는 해당 항의 t 통계량 값이 3 이상의 큰 값을 보이는 것으로도 확인된다. 그리고 회귀분석 모형의 설명력을 나타내는 조정된 결정계수는 0.264를 기록하였다. 이러한 결과는 표본 기업들에 있어서 디지털 공급사슬관리의 실행은 성과에 정(+)의 영향력을 지닌다는 의미를 제공한다. 이 기본 모형에서의 조정된 결정계수의 값이 0.3 이하를 나타내고 있는데 다수의 기존 연구들에서 그러한 경우에도 모형의 설명력이 인정되고 있으며 최근의 연구의 예로 Hong et al.(2020)을 들 수 있다.

본 연구의 두 번째 연구가설은 디지털 공급사슬관리 실행과 성과 간의 관계에 대하여 AI의 적극적 활용이 정(+)의 조절 효과를 나타내는지 검증하는 가설이었다. 이의 검증을 위하여 조절적 혹은 위계적 회귀분석 모형에 의한 검증이 실행되었으며, 그 결과는 아래 <표 4>의 Step 2에 제시되어 있다. 즉 디지털 공급사슬관리 실행 그리고 디지털 공급사슬관리와 AI의 적극적 활용의 곱에 해당하는 변수를 독립변수로 하고 기업 성과를 종속변수로 하는 다중 회귀분석을 수행하였다.

위 <Table 4>의 Step 2에 보이는 결과와 같이 회귀분석 모형 전체의 유의성을 나타내는 F 통계량 값은 16.109로서 1% 유의수준에서 유의한 결과를 나타내었고, 1.2986을 나타낸 상수항 그리고 0.457인 디지털 공급사슬관리 실행 변수의 계수는 1% 유의수준에서 유의한 결과를 보였으나, 디지털 공급사슬관리 변수와 AI 활용 변수를 곱한 값에 해당하는 변수의 계수는 0.060으로 유의하지 않은 결과를 보였다. 모형의 설명력을 의미하는 조정된 결정계수는 0.294로서 Step 1의 회귀모형보다 0.004 만큼만 증가하는데 그치는 결과를 보였다. 이와 같은 실증 분석 결과를 종합하면 두 번째 연구가설은 한정적으로만 채택할 수 있음을 알 수 있다.

연구의 세 번째 단계로서, 디지털 공급사슬관리 실행과 성과와의 관계에 대해 빅데이터 활용이 조절적 효과를 나타내는지 검증하는 것은 앞 장에서 설정된 세 번째 연구가설이며, 역시 이 가설도 조절적 혹은 위계적 회귀분석 모형에 의해 검증되었다. 즉 디지털 공급사슬관리 실행 그리고 디지털 공급사슬관리와 빅데이터의 적극적 활용의 곱에 해당하는 변수를 독립변수로 하고 기업 성과를 종속변수로 하는 다중 회귀분석을 수행하였다.

그 결과는 아래의 <Table 5>의 Step 3에 제시하였다. 모형의 유의성을 나타내는 F 통계량 값은 7.184로서 Step 1보다 대폭 감소하였지만 여전히 1% 유의수준에서 유의한 결과를 보였으며, 1.302의 상수항, 0.481의 디지털 공급사슬관리 실행 변수의 계수 그리고 0.100을 기록한 디지털 공급사슬관리 실행 변수와 빅데이터 활용 변수를 곱한 값에 해당하는 변수의 계수도 모두 1% 유의수준에서 유의한 결과를 나타내고 있다. 설명력 측면에서도 조정된 결정계수가 0.277을 기록하여 Step 1의 모형보다 0.013 증가하여 연구가설 2보다 더 증가했음을 보여 주고 있다. 이와 같은 실증 분석 결과들을 종합하면 세 번째 연구가설은 성공적으로 채택되었다.

이러한 실증 분석 결과의 시사점은 중국에 소재한 표본 업체들에 있어서 디지털 공급사슬관리 실행과 빅데이터 분석의 적극적 활용은 시너지적 효과를 나타내어 보다 나은 성과 개선을 가능하게 하지만, AI의 적극적 활용은 한정적으로만 그러한 효과를 나타낸다는 것이다. 전자는 검정 통계량과 계수의 유의성 면에서 모두 유의한 결과를 나타냈으나 후자의 경우에는 검정통계량은 유의하지만 조절변수에 해당하는 계수가 유의하지 않은 결과를 나타내었다.

이러한 결과의 원인은 빅데이터 분석은 자료의 분석

Table 4. Relationship between Digital Supply Chain Management and Performance and Moderating Effect of AI Usage

Variable	Step 1			Step 2		
	B	Standard deviation	t value	B	Standard deviation	t value
Constant	1.523**	0.166	9.164	1.298**	0.172	7.543
Digital Supplier Chain Mgt.	0.520**	0.046	11.281	0.457**	0.048	9.583
DSCM * AI Usage	-	-	-	0.060	0.052	1.141
R ²	0.266			0.298		
adj-R ²	0.264			0.294		
adj-ΔR ²	-			0.004		
F	127.269**			16.109**		

** p < 0.01.

Table 5. Relationship between Digital Supply Chain Management and Performance, and Moderating Effect of Big Data Usage

Variable	Step 1			Step 3		
	B	Standard deviation	t value	B	Standard deviation	t value
Constant	1.523**	0.166	9.164	1.302**	0.184	7.074
Digital Supplier Chain Mgt.	0.520**	0.046	11.281	0.481**	0.048	10.039
DSCM * Big Data Usage	-	-	-	0.100**	0.037	2.680
R ²	0.266			0.281		
adj-R ²	0.264			0.277		
adj-ΔR ²	-			0.013		
F	127.269**			7.184**		

**p < 0.01.

과 해석을 통하여 DSC과 연계되어 즉각적으로 성과에 긍정적 영향을 미치지만, AI의 경우 범위가 넓기 때문에 장기적으로는 성과에 긍정적 영향을 제공할 것으로 예상되지만 단기적으로는 성과의 향상에 영향을 미치지 않는다는 해석을 가능하게 한다. 이는 AI의 경우 채택 및 활용에서 성과의 향상까지 상당한 시간이 경과되어야 한다는 ‘정보기술에서의 생산성 역설’과 어느 정도 연관되는 결과로 파악될 수 있다(Brynjolfsson, 1993).

5. 결 론

본 연구는 중국 기업들을 연구 표본으로 하여 디지털 공급사슬관리 실행이 성과에 정(+)의 영향력을 제공하는지를 검증하는 것을 첫 단계의 목표로 설정하였다. 연구의 다음 단계에서는 디지털 공급사슬관리 실행과 성과 간의 관계에 대해 최근 관심이 집중되고 있는 4차 산업혁명의 주요 구성요소 및 디지털 공급사슬관리의 기반 기술들 중 AI의 활용과 빅데이터 분석의 활용이 각기 정(+)의 조절 효과를 나타내는지 검증하고자 하였다.

중국 내에 소재한 353개 업체들을 대상으로 하여 설문지에 의한 서베이 조사가 수행되었고 그 결과는 회귀분석 및 조절적 회귀분석 모형에 의해 실증적으로 검증이 이루어졌다. 실증적 분석 결과, 첫 번째로 중국 내 표본 기업들의 디지털 공급사슬관리 실행은 성과에 대하여 정(+)의 영향을 제공한다는 가설은 성공적으로 채택되었다.

그 다음 단계의 가설 검증을 위하여 AI의 활용과 빅데이터 분석의 활용을 각기 조절 변수로 한 위계적 회귀분석 모형에 의한 검증을 실행하였고 그 결과로 AI 활용은 한정적 조절 효과를 빅데이터 분석은 정(+)의 조절 효과를 갖는 것으로 나타났다. 이는 중국 내 기업들에 있어서 디지털 공급사슬관리 실행과 빅데이터 분석의 활용은

서로 적합성을 가지면서 성과를 보다 개선시키는 효과를 갖지만 AI의 활용은 한정적으로만 그러한 효과를 나타낸다는 것을 의미한다. 이러한 실증 분석 결과는 빅데이터 분석의 경우 자료의 처리와 해석에 의하여 디지털 공급사슬관리와 결합하여 즉각적으로 성과 개선에 긍정적 효과를 갖지만, AI 활용은 광범위한 기술로서 특성을 지니고 있기 때문에 디지털 공급사슬관리와 연결되어 성과 개선에 공헌하기 위해서는 시간의 경과가 필요하다는 해석을 가능하게 하는데, 이는 정보기술들은 성과 개선에 공헌하기 위하여 시간의 경과가 필요하다는 고전적인 정보기술의 ‘생산성 역설’(Brynjolfsson, 1993)의 관점과 연결되는 것으로 파악될 수 있다.

본 연구의 이론적 공헌으로는 디지털 공급사슬관리의 개념을 정리하고 그것이 성과에 정(+)의 영향을 미치는지를 그리고 4차 산업혁명의 다양한 구성요소 및 기술들이 디지털 공급사슬관리와 연계되어 성과의 향상에 공헌하는지를 실증적으로 검증하였다는 점을 들 수 있을 것이다. 실무적 관점에서는 본 연구는 디지털 공급사슬관리 실행과 빅데이터 활용의 연계를 통한 더 나은 성과 향상이 가능하다는 그리고 광범위한 기술로서의 AI 활용이 성과 개선에 공헌하기 위해서는 시간의 경과가 요구된다는 시사점을 제공하였다는 점이다. 향후에는 4차 산업혁명의 다른 기술들과 디지털 공급사슬관리 및 다른 새로운 접근과의 연계에 대한 연구가 새로운 주제로서 추진될 수 있을 것이다.

REFERENCES

- [1] Acimovic, S. & Stajic, N. (2019). Digital supply chain: Leading technologies and their impact on industry 4.0. *Proceedings of 19th International scientific conference:*

- Business Logistics in Modern Management*, October, Osijek, Croatia, 75-90.
- [2] Bar, K., Herbert-Hansen, Z. and Khalid, W. (2018). Considering industry 4.0 aspects in the supply chain for an SME. *Production Engineering*, 12, 747-758.
- [3] Bhargava, B., Ranchal, R., & BenOthmane, L. (2013). Secure information sharing in digital supply chains. *Proceedings of 3rd IEE International Advanced Computer Conference*, 1636-1640.
- [4] Brynjolfsson, E.(1993). The productivity paradox of information technology. *Communication of the ACM*, 36(12), 67-77.
- [5] Buyukozkan, G. & Gocer, F. (2018). Digital Supply Chain: Literature review and a proposed framework for future research. *Computers in Industry*, 97, 157-177.
- [6] Garay-Rondero, C., Martinez-Flores, J., Smith, N., Morales, S., & Aldrette-Malcara, A.(2020). Digital supply chain model in Industry 4.0, *Journal of Manufacturing Technology Management*, 31(5), 887-933.
- [7] Hong, S., Bauer, J., Lee, K., & Granados, N. (2020). Drivers of supplier participation in ride-hailing platforms, *Journal of Information Systems*, 37(3), 602-630.
- [8] Kinnet, J.(2015). Creating digital supply chain: Monsanto's journey, *SlideShare*, 1-16.
- [9] Lee, J. H.(2018). *The Fourth Industrial Revolution and New Growing Power in the Future*, Jinhan M&B.
- [10] Li, G., Yang, H., Sun, L., & Sohal, A. (2009). The impact of IT implementation on supply integration and performance. *International Journal of Production Economics*, 120, 125-138.
- [11] Linh, N., Kumar, V., & Ruan, X. (2019). Exploring enablers, barriers and opportunities to digital supply chain management in vietnamese manufacturing SMEs. *International Journal of Organizational Business Excellence*, 2(2), 101-120.
- [12] Liu, H., Huang, Q., & Wei, S.(2015). The impacts of IT capability on internet-enabled supply and demand process integration, and firm performance in manufacturing and services. *The International Journal of Logistics Management*, 26(1), 172-194.
- [13] Liu, H., Ke, W., Wei, K., & Hua, Z.(2013). The impact of IT capability on firm performance: The mediating roles of absorptive capacity and supply chain agility, *Decision Support System*, 54, 1452-1462.
- [14] Rai, A., Patnayakuni, R., & Seth, N.(2006). Firm performance impacts of digitally enabled supply chain integration capabilities. *MIS Quarterly*, 30(2), 225-246.
- [15] Wong, C., Lai, K., & Cheng, T.(2011). Value of information integration to supply chain management: Roles of internal and external contingencies. *Journal of Management Information Systems*, 28(3), pp. 161-199.
- [16] Xue, L., Zhang, C., Ling, H., & Zhao, X. (2013). Risk mitigation in supply chain digitization: System modularity and information technology governance. *Journal of Management Information Systems*, 30(1), 325-352.



김 영 길

서울대학교 경영학 박사
 현재: 신한대학교 글로벌통상경영학과
 조교수
 관심분야: SCM, 서비스 경영, 중국 경영
 등



박 정 수

연세대학교 정치학사
 서울대학교 경영학 석사, 박사
 현재: 중앙대학교 다빈치 교양대학
 조교수
 관심분야: SCM, 서비스경영, OR응용 등

혼성제조시스템을 위한 생산용량할당정책 수립과 지연재고정책의 이익 성과 효과성 분석*

김은갑**†

**이화여자대학교, 경영대학

Capacity Allocation for a Hybrid MTS-MTO Manufacturing System and Effectiveness Analysis of Backlogging for Profit Performance*

Eungab Kim**

**College of Business Administration, Ewha Womans University

We consider a hybrid manufacturing system that produces components based on both long-term contract with higher priority and short-term contract with lower priority. The manufacturer operates production on a make-to-stock mode for the long-term contract and on a make-to-order mode for the short-term contract. Recently, hybrid manufacturing has received attention from academia and industry because it can contribute to profit enhancement through revenue channel diversification. In this paper, using a Markov decision process model, we study the structure of capacity rationing and admission control policies under the assumption of non-preemptive batch production and compound Poisson demand process. We also investigate the effectiveness of backlogging short-term contract orders on profit and examine the impact of long-term contract demand sizes on profit with varying probabilities of each demand size and variance of demand.

Keyword : Admission control, Capacity Allocation, Hybrid manufacturing, Compound Poisson Process, Markov decision process

* 이 논문은 2019년 대한민국 교육부와 한국연구재단의 인문사회분야 중견연구자지원사업의 지원을 받아 수행된 연구임(NRF-2019S1A5A2A01044200).

† **Corresponding Author** : College of Business Administration, Ewha Womans University, 52, Ewhayeodae-gil, Seodaemun-gu, Seoul, 03760, Korea.
Tel: +82-2-3277-3970, E-mail: evanston@ewha.ac.kr

Received: 9 November 2021, **Revised**: 7 December 2021, **Accepted**: 12 December 2021

1. 서 론

본 논문은 동일 생산시설에서 장기계약 대응을 위한 계획생산(make-to-stock, 이하 MTS)과 단기계약 대응을 위한 주문생산(make-to-order, 이하 MTO)을 함께 수행하고 있는 부품제조기업을 대상으로 한다. 해당 생산시설은 혼성(Hybrid) 생산시스템으로 시점 별로 한 가지 생산만 가능하기 때문에 생산 용량을 MTS와 MTO에 분배해야 하는 문제와 단기계약주문을 통제해야 하는 문제에 직면하게 된다. <Fig. 1>은 본 논문의 모형을 도식화한 것이다. 제조업체의 한정된 생산용량을 재고보충에 과도하게 분배한다면 단기계약주문 대응에 필요한 생산용량 확보가 어려워질 수 있다. 반대로 단기계약주문 수용 강화 시, 장기계약 이행에 필요한 재고 확보에 차질이 생길 수 있다. 따라서 제조업체의 이익 고도화를 위해서는 생산용량할당과 주문통제 정책이 통합적으로 이루어져야 할 필요성이 있다. 본 논문은 최적 생산용량할당과 단기계약주문통제 정책을 분석하고, 단기계약주문에 대한 지연재고 허용이 이익 성과 측면에서 효과적인 정책이 될 수 있는지를 검증한다.

전통적으로 부품제조업체들은 완제품제조업체와의 장기공급계약을 통해 수익을 확보해왔다. 그러나 판매후 서비스시장의 성장(Cohen et al., 2006)과 부품-완제품 두 단계 공급망에서 부품시장이 주요 수익원으로 등장(Iravani et al., 2012)하면서 부품업체들의 생산용량 운용이 중요하게 되었다. 그러나 장기계약수요와는 달리 단기계약수요는 예측이 어렵기 때문에, 효과적인 생산용량 운영은 힘든 과제가 될 수 있다. Bish et al.(2005)은 주문생산시스템에서 지연재고수준에 따라 생산용량을 할당하는 것이 주문수요변동성 완화에 효과적임을 보였다. 본 논문에서도 생산용량을 장기계약수요와 단기계약주문에 동태적으로 배분하는 정책을 제시한다.

최근 맞춤형제품생산이 주목을 받으면서 주문통제와 생산일정수립에 대한 연구가 활발히 진행되고 있다. 주문통제의 핵심은 수용/거절 판단과 지연재고 허용이며, 수용여부는 지연재고수준에 의해 결정된다. 본 논문과

동일 모형을 연구한 Carr and Duenyas(2000)과 Iravani et al.(2012)는 단기계약주문에 대해 지연재고를 허용한다. 여기서 제기되는 의문은 ‘단기계약주문에 대해 지연재고를 허용하는 것이 이익 성과 측면에서 바람직한가?’이다. 본 논문에서는 단기계약주문에 대해서 지연재고를 허용하는 시스템과 생산유휴 시 발생한 단기계약주문에 한해 수용여부를 결정하는 시스템 간 이익 수행도 비교를 통해 여기에 대한 답을 제시하고자 한다.

2. 문헌 연구

본 논문과 밀접한 연구 분야는 한 단계 MTS-MTO 혼성제조시스템이다. Carr and Duenyas(2000)와 Iravani et al.(2012)는 MTO생산보다 MTS생산이 우선순위일 때, 생산용량배분과 주문통제 정책을 연구하였다. 두 선행 연구들은 생산용량배분과 주문통제 정책을 연구했다는 점에서 본 논문과 매우 유사하나 세 가지 측면에서 차이가 있다. 첫째, 두 선행 연구들에서 MTS수요 크기는 항상 1인 반면, 본 논문에서 MTS수요 크기는 가변적이다. Elhafsi(2009)는 고객수요의 크기가 확률적으로 주어지는 상황을 복합포아송분포를 사용하여 모형화 할 수 있음을 보이고 있다. 둘째, 두 선행 연구들에서 생산 Lot 크기는 1이고 고정비용은 고려하지 않았다. 반면, 본 논문에서 생산 Lot 크기는 배치이고, 고정비용을 고려한다. 셋째, 두 선행 연구들은 중단 가능(preemption)한 생산프로세스를 가정하고 있다. 이에 반해, 본 논문은 종료 전에는 중단 불가(non-preemption)인 생산프로세스를 가정한다. Preemption을 가정할 경우, 마코프의사결정모형에서는 주문수용 또는 수요충족 시, 진행 중인 생산을 중단하고 상태 전이가 이루어진다. 그러나 마코프의사결정모형의 memoryless 특성 때문에 다시 생산이 진행되는 상황이 계속해서 발생할 수 있다.

Wang et al.(2019)은 MTS제품생산일정 수립 후, MTO제품 주문 발생 시 해당 주문생산 가능 여부에 따라 수용되는 문제를 분석하였다. Wang et al.(2019)의 모형은 본 논문의 지연재고 정책 효과성 분석에서 사용될 모형의

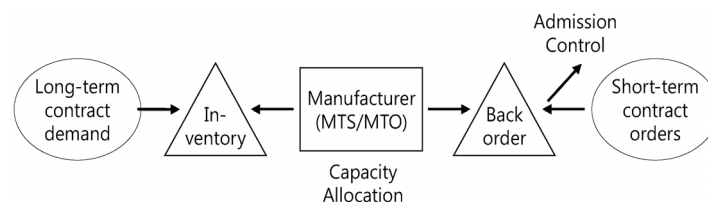


Fig. 1 A Hybrid Manufacturing System with Long-Term and Short-term Contracts

토대이다. Abedi and Zhu(2020)는 고정 크기의 MTS수요에 대해 전부 수용, 부분 수용, 또는 수용거절 문제를 연구하였다. Kuthambalayan and Bera(2020)는 MTS제품과 MTO제품이 동일한 반제품을 사용할 때, 생산용량배분, MTS제품재고수준, 반제품재고수준을 결정하는 문제를 연구하였다. Peeters and van Ooijen(2020)는 혼성제조시스템에 대한 문헌 연구를 진행하였다.

한 단계 MTS-MTO 혼성제조시스템은 한 단계-복수고객 제조시스템의 특수한 경우로 볼 수 있다. MTO생산에서 지연재고크기에 제한이 없고 전환비용이 없다면, $c\mu$ 규칙이 최적이다(Buyukkoc et. al, 1985). 전환비용이 있는 경우, 계층지연재고비용이 동일한 시스템은 Hofri and Ross(1987)와 Liu et al.(1992), 상이한 시스템은 Duenyas and Van Oyen(1996)에 의해 연구되었다. Kim and Van Oyen(2021)은 계층지연재고비용이 상이할 때 최적 생산용량할당과 주문통제 정책을 증명하였다. MTS생산에서 재고보충과 수요통제는 Ha(1997a)에 의해서 처음 연구되었다. Ha(1997a)는 거절된 수요는 판매손실로 가정하였다. Ha(1997b)와 de Vericourt et al.(2002)는 거절된 수요가 지연재고로 처리되는 상황을 연구하였다. Ding et al.(2016)은 생산고정비용이 발생할 때, 재고 할당과 지연재고 문제를 연구하였다. Kouki et al.(2020)은 계층 임계치를 설정하고, 재고수준이 계층의 임계치를 초과하는 경우에 한해 수요를 충족시키는 정책을 제시하였다.

3. 주요 가정 및 마코프 의사결정 모형 수립

본 논문의 주요 가정은 다음과 같다:

- 생산시간은 평균이 μ^{-1} 인 확률 프로세스를 따른다. 생산은 크기가 q 인 배치 방식이며, 고정 비용이 발생한다.
- 장기계약수요(이하 LD)는 단위 시간당 평균발생빈도가 λ_1 이며, 크기는 최소 1, 최대 D 인 확률변수이다. 미 충족된 LD는 지연재고로 처리된다.
- 단기계약주문(이하 SO)은 단위 시간당 평균발생빈도가 λ_2 이며, 주문량의 크기는 배치 생산량과 동일하다. SO 수용거절 시, 거절비용은 고려하지 않는다. 생산된 물량은 즉시 출고된다.
- Carr and Duenyas(2000)와 Iravani et al.(2012)에서는 지연 중인 SO에 대해 비용 발생을 가정하고 있다. 본 논문도 동일한 가정을 한다.

논문에서 사용될 주요 용어들을 정리하면 <Table 1>과 같다.

Table 1. Description of Model Parameters

Notation	Description
$R_1(R_2)$	Unit revenue of MTS(MTO) product
D	Maximum size of MTS demand
λ_1	MTS demand rate
$p(i)$	Probability that MTS demand size is i
λ_2	MTO demand rate
μ^{-1}	Mean of production time
c_K	Setup cost of production
$b_1(b_2)$	Backlog cost rate of unit MTS(MTO) demand
c_h	Holding cost rate of unit MTS product
q	Batch size of production
x_1	Inventory level of MTS product
x_2	Backorder level of MTO demand
n	No production in process($n=0$); MTS in process($n=1$); MTO in process($n=2$)
(x, n)	System state where $x=(x_1, x_2)$
e_l	Two-dimensional l^{th} unit vector
γ	Transition rate, $\gamma \equiv \lambda_1 + \lambda_2 + \mu$
g	Optimal average profit during γ^{-1}

본 논문에서는 마코프의사결정(이하 MDP)모형을 이용하여 최적 정책의 구조를 분석한다. 이를 위해 생산시간은 지수분포, LD발생은 복합포아송분포, SO발생은 포아송분포로 가정한다. MDP모형의 상태 전이는 다음과 같다: $(x, 0)(x_2 > 0)$ 에서 LD생산결정 시 $(x, 1)$, SO생산결정 시 $(x, 2)$ 로 전이; $(x, 0)(x_2 = 0)$ 에서 LD생산결정 시 $(x, 1)$, 생산유휴결정 시 $(x, 0)$ 으로 전이; (x, n) 에서 크기가 i 인 LD 발생 시 $(x - ie_1, n)$ 으로 전이; (x, n) 에서 SO 발생 시, 수용하면 $(x + e_2, n)$, 거절하면 (x, n) 으로 전이; $(x, 1)$ 에서 생산완료 시 $(x + qe_1, n)$ 으로 전이; $(x, 2)$ 에서 생산완료 시 $(x - e_2, n)$ 으로 전이.

본 논문의 벨만방정식은 가치반복수식자 T 를 통해서 정의된다(Puterman, 2005). $H(x) \equiv x_1 c_h 1(x_1 > 0) - x_1 b_1 1(x_1 < 0) + x_2 b_2$ 라고 두자. $H(x)$ 는 지연재고비용과 재고유지비용의 합이며, $1(a)$ 은 a 가 성립하면 1, 성립하지 않으면 0이 된다. $f(x, n)$ 을 (x, n) 에서의 가치함수라고 두면, T 는 다음과 같다:

$$\begin{aligned}
 & Tf(x, n) \\
 &= \begin{cases} \max\{T^R f(x, n), T^E f(x, n)\}, & n=0, x_2 > 0 \\ \max\{T^N f(x, n), T^R f(x, n)\}, & n=0, x_2 = 0 \\ T^N f(x, n), & n=1 \end{cases} \quad (1)
 \end{aligned}$$

여기서

$$T^N f(x, n) = \frac{1}{\gamma} [-H(x) + \sum_{k=1}^2 \lambda_k T_k f(x, n) + \mu T_3 f(x, n)] \quad (2)$$

$$T^R f(x, 0) = -c_K + T^N f(x, 1) \quad (3)$$

$$T^E f(x, 0) = -c_K + T^N f(x, 2) \quad (4)$$

여기서

$$T_1 f(x, n) = \sum_{i=1}^D p(i) (iR_1 + f(x - ie_1, n)) \quad (5)$$

$$T_2 f(x, n) = \max \{f(x + e_2, n) + qR_2, f(x, n)\} \quad (6)$$

$$T_3 f(x, n) = \begin{cases} f(x, 0), & n = 0 \\ f(x + qe_2, 0), & n = 1 \\ f(x - e_2, 0), & n = 2 \end{cases} \quad (7)$$

식 (1)은 생산용량할당 의사결정을 나타낸다. $n=0$ 이고 $x_2 > 0$ 이면 LD생산(T^R)과 SO생산(T^E), $n=0$ 이고 $x_2=0$ 이면 LD생산(T^R)과 생산유휴(T^N)에 대한 의사결정이 요구된다. λ_1 항목과 T_2 는 LD 발생 직후의 전이를 나타낸다. 미 충족된 수요는 지연재고가 된다. λ_2 항목과 T_3 는 SO 수용/거절 의사결정과 전이를 나타낸다. μ 항목과 T_4 는 생산완료 직후의 전이를 나타낸다. 식 (7)에서 $n=0$ 이면 생산유휴이기 때문에 상태변화가 없다.

본 논문의 단위 시간당 평균이익기준 MDP 모형의 벨만방정식은 다음과 같다:

$$g + f(x, n) = Tf(x, n) \quad (8)$$

의사결정변수인 g 는 가치반복알고리즘을 통해 구하게 된다. 모든 (x, n) 에 대해서 알고리즘의 k 번째 단계에서 구한 $f^k(x, n)$ 과 가치반복수식자 T 를 적용한 $Tf^k(x, n)$ 간의 편차가 알고리즘 종료 기준보다 작게 될 때까지 T 를 반복적으로 적용시켜 g 를 구하게 된다. 식 (1)-(8)로 정의된 MDP 모형을 모형 I으로 표기하기로 한다.

4. 최적 정책의 구조

최적 생산용량할당과 SO통제정책은 다음 함수식들로 정의 된다:

$$P(x_2, 0) := \max \{x_1 : T^O f(x, 0) \geq T^E f(x, 0)\} \quad (9)$$

$$C(0, 0) := \max \{x_1 : T^N f(x, 0) \leq T^O f(x, n)\} \quad (10)$$

$$A(x_1, n) := \max \{x_2 : f(x + e_2, n) + qR_2 \geq f(x, n)\} \quad (11)$$

$C(0, 0)$ 은 LD생산이 생산유휴보다 이익인 x_1 중에서 가장 큰 값이다. $P(x_2, 0)$ 는 지연재고수준이 x_2 일 때, LD생산이 SO생산보다 이익인 x_1 중에서 가장 큰 값이다. $A(x_1, n)$ 은 신규SO 수용이 거절보다 이익인 x_2 중에서 가장 큰 값이다.

<Fig. 2>는 종료조건 $\epsilon = 0.001$ 을 갖는 가치반복알고리즘을 통해 구한 최적 정책의 구조를 제시하고 있다. 사용된 예제는 다음과 같다: $R_1 = 15$, $R_2 = 16.5$, $\lambda_1 = 4$, $D = 4$, $p(i) = 0.25$, $i = 1, 2, 3, 4$, $\lambda_2 = 0.35$, $\mu^{-1} = 1$, $q = 25$, $c_K = 200$, $c_h = 1$, $b_1 = 5$, $b_2 = 15$. 모수 값들은 Carr and Duenyas(2000)와 Iravani et al.(2012), 그리고 Elhafsi (2009)의 수치실험에서 사용된 값들을 참조하였다.

<Fig 2>에서 x_1 축은 부품재고수준, x_2 축은 지연재고수준을 나타낸다. $C(0, 0)$ 는 지연재고가 없을 때 LD생산 결정 임계 값이다. 본 예제에서는 $C(0, 0) = 15$ 이다. 따라서 $x_1 \leq 15$ 이면 LD생산을 진행한다. $P(x_2, 0)$ 는 최적 생산용량할당 정책이다. $x_1 \leq P(x_2, 0)$ 이면 LD생산, $x_1 > P(x_2, 0)$ 이면 SO생산을 진행한다. $A(x_1, 0)$, $A(x_1, 1)$, $A(x_1, 2)$ 는 부품재고수준이 x_1 일 때, 각각 $n=0$, $n=1$, $n=2$ 에 대응하는 최적 SO통제정책이다. $C(0, 0)$, $P(x_2, 0)$, $A(x_1, n)$ 은 상태공간을 세부 영역으로 나눈다. 예를 들면, (17,7,0)에서는 SO생산을 진행하고, 신규SO는 거절한다. (10,2,0)에서는 LD생산을 진행하고, 신규SO는 수용한다.

Carr and Duenyas(2000)와 Iravani et al.(2012)와 비교했을 때, <Fig. 2>에 제시된 최적 정책의 가장 큰 차이점은 SO 통제에 있다. Carr and Duenyas(2000)와 Iravani et al.(2012)은 Preemption을 가정하기 때문에 SO통제정책은 x_1 과 x_2 에 의해서만 영향을 받는다. 이에 반해, 본 논문은 Non-preemption을 가정하기 때문에 SO수용여부는 생산 상태에 의해서도 영향을 받는다. 따라서 x_1 과 x_2 가 동일해도 n 값에 따라 수용여부가 달라질 수 있다. <Fig. 2>는 $A(x_1, 1) \leq A(x_1, 0) \leq A(x_1, 2)$ 임을 보여준다. 예를 들면, SO 발생 시, (15,4,2)에서는 수용하지만, (15,4,0)과 (15,4,1)에서는 거절된다. (16,4,0)과 (16,4,2)에서는 신규SO가 수용되지만, (16,4,1)에서는 거절된다. $n=2$ 일 때는 생산완료 시 x_2 가 1 감소하게 되어 주문수용에 따른 지연재고부담이 적기 때문인 것으로 이해할 수 있다. 반면, $n=1$ 일 때는 재고보충생산이 진행 중이므로 신규SO 수용에 따른 지연재고부담이 가장 커서

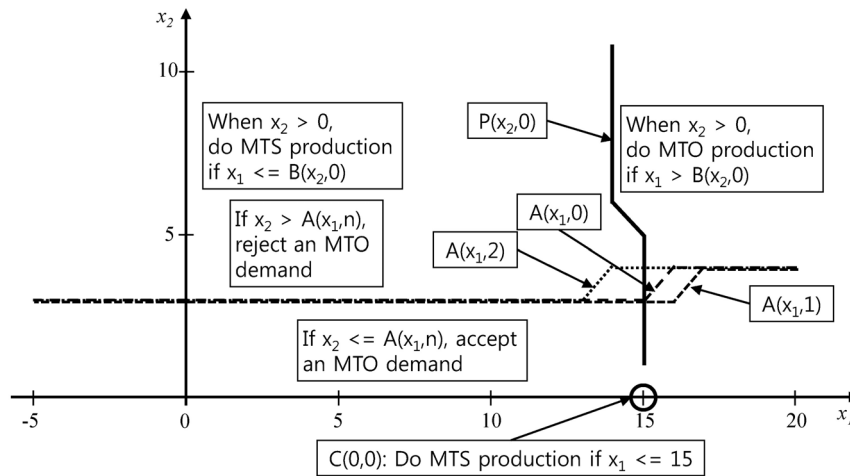


Fig. 2 The Structure of Optimal Policies

$n=0$ 과 $n=2$ 일 때보다는 SO 수용이 제한적으로 이루어져야 하는 것으로 이해할 수 있다.

한편, <Fig. 2>는 최적 생산용량할당 정책인 $P(x_2,0)$ 가 x_2 에 대해 감소함수임을 보여준다. 이 결과는 지연재고 수준이 증가함에 따라 LD 생산보다는 SO 생산의 비중이 확대되어야 함을 의미한다.

5. 단기계약주문의 지연재고정책 효과성

이번 절에서는 SO 에 대한 지연재고정책의 효과성을 분석한다. 이를 위해, SO 발생 시, $n=0$ 일 때 한해 수용 여부를 결정하는 모형(이하 모형 II)을 제시한다. 모형 I과 비교했을 때, 모형 II의 수식자 T_2 와 T_3 는 다음과 같이 다르다:

$$T_2f(x,n) = \begin{cases} \max\{f(x+e_2,2) + qR_2, f(x,0)\}, & n=0 \\ f(x,n), & n>0 \end{cases} \quad (12)$$

$$T_3f(x,n) = \begin{cases} f(x,0), & n=0 \\ f(x+qe_2,0), & n=1 \\ f(x,0), & n=2 \end{cases} \quad (13)$$

식 (12)는 (x,n) 에서 SO 발생 시, $n>0$ 이면 거절된다. $n=0$ 일 때, 수용되면 n 은 2의 값을 갖게 된다. 식 (13)은 생산종료 시, n 값에 따른 전이를 나타낸다. $n=2$ 이면, 생산물량은 바로 출고되기 때문에 전이 전후 상태는 동일하다.

<Table 2>는 지연재고비용 관점에서 모형 I과 II의 수행도를 비교한다. b_1 과 b_2 값 설정은 Carr and Duenyas(2000)와 Iravani et al.(2012)을 참고하였다. 이 논문들에서 매출 단가 대비 지연재고비용의 비율 최소값은 0.02이고 최대값은 0.2였다. 본 논문에서는 b_2/qR_2 의 최소값은 0.07, 최대값은 0.17로 설정하였다. 장기계약이행 시 재고부족은 바람직하지 않기 때문에 $b_1/R_1 \geq 0.5$ 로 설정하였다. %열에 있는 숫자들은 모형 I과 II 간 수행도 차이를 백분율로 나타낸 것이다: $(g^I - g^{II})/g^{II} \times 100$. SO 에 대한 지연재고정책의 효과성은 b_2/qR_2 값에 의해 극명하게 달라짐을

Table 2. Performance Comparison between Model I and II Varying b_1 and b_2

			$b_1=10$		$b_1=15$		$b_1=20$		$b_1=25$		$b_1=30$	
				%		%		%		%		%
b_2	b_2/qR_2	g^{II}	66.0		63.9		62.3		60.9		59.8	
35	0.07	g^I	87.9	33.2	85.7	34.1	84.0	34.9	82.6	35.5	81.4	36.1
45	0.09	g^I	77.5	17.4	75.3	17.9	73.6	18.2	72.2	18.6	71.0	18.8
55	0.11	g^I	68.5	3.8	66.3	3.8	64.6	3.8	63.2	3.8	62.0	3.8
58	0.116	g^I	66.1	0.2	64.0	0.1	62.2	0.0	60.8	-0.1	59.6	-0.2
65	0.13	g^I	61.0	-7.6	58.8	-7.9	57.1	-8.2	55.7	-8.5	54.5	-8.8
75	0.15	g^I	53.9	-18.3	51.8	-19.0	50.1	-19.6	48.7	-20.1	47.5	-20.6
85	0.17	g^I	47.1	-28.7	44.9	-29.7	43.3	-30.5	41.9	-31.3	40.7	-32.0

알 수 있다. $b_2/qR_2 \leq 0.11$ 이면 모형 I이 II보다 우수하지만 $b_2/qR_2 \geq 0.13$ 이면 결과는 반대가 된다. $b_2/qR_2 = 0.116$ 이면 모형 I과 II 간 수행도 차이는 거의 발생하지 않았다.

다음은 LD 발생 확률 $p(i)$ 가 모형 I과 II의 이익 수행도에 미치는 영향을 분석한다. 본 수치실험에서는 $D=5$ 일 때, 총 5가지 조합의 $(p(1), \dots, p(5))$ 를 고려하였다 (<Table 3> 참조). <Table 4>는 <Table 3>의 예제 1-5에 대해서 구한 모형 I과 II의 수행도를 제시한다. <Table 4>에서 고려한 b_1 과 b_2 의 조합은 $b_1 \in \{10, 15, 20, 25, 30\}$ 과

$b_2 \in \{35, 75\}$ 이다. <Table 4>를 보면, 예제 1($p(i)$ 가 i 에 대해서 증가함수)에서 모형 I과 II의 수행도 차이는 확대되고, 예제 5($p(i)$ 가 i 에 대해서 감소함수)에서는 축소됨을 알 수 있다. 한편, 예제 2-4가 모형 I과 II의 수행도 차이에 미치는 영향은 거의 동일함을 알 수 있다. 예제 2-4에서 나타나는 공통적인 특징은 $p(1) = p(5)$ 이다. 이 결과를 토대로 $p(1)$ 과 $p(5)$ 의 비대칭이 커질수록 모형 I과 II의 수행도 차이에 미치는 영향은 확대될 것으로 유추해볼 수 있다.

Table 3. Setting of $p(1), \dots, p(5)$

Ex	$(p(1), p(2), p(3), p(4), p(5))$	Description
1	(0.05, 0.1, 0.15, 0.2, 0.5)	Increasing in i
2	(0.3, 0.15, 0.1, 0.15, 0.3)	Reverse symmetric triangular
3	(0.2, 0.2, 0.2, 0.2, 0.2)	Constant
4	(0.1, 0.15, 0.5, 0.15, 0.1)	Symmetric triangular
5	(0.5, 0.2, 0.15, 0.1, 0.05)	Decreasing in i

Table 4. Comparison of g^I and g^{II} using Examples 1-5 in <Table 3>

Ex			$b_1=10$		$b_1=15$		$b_1=20$		$b_1=25$		$b_1=30$	
				%		%		%		%		%
1		g^{II}	50.3		47.5		45.3		43.4		41.9	
	$b_2=35$	g^I	71.1	41.4	67.6	42.5	64.7	43.0	62.1	43.0	59.8	42.8
	$b_2=75$	g^I	39.4	-21.6	36.5	-23.1	34.2	-24.5	32.1	-26.2	30.1	-28.1
2		g^{II}	65.7		63.6		62.0		60.6		59.4	
	$b_2=35$	g^I	87.6	33.3	85.4	34.2	83.6	35.0	82.2	35.7	81.0	36.3
	$b_2=75$	g^I	53.6	-18.4	51.4	-19.1	49.7	-19.7	48.3	-20.3	47.1	-20.7
3		g^{II}	66.0		63.9		62.3		60.9		59.8	
	$b_2=35$	g^I	87.9	33.2	85.7	34.1	84.0	34.9	82.6	35.5	81.4	36.1
	$b_2=75$	g^I	53.9	-18.3	51.8	-19.0	50.1	-19.6	48.7	-20.1	47.5	-20.6
4		g^{II}	66.4		64.3		62.7		61.4		60.2	
	$b_2=35$	g^I	88.3	33.1	86.1	33.9	84.4	34.7	83.1	35.4	81.9	35.9
	$b_2=75$	g^I	54.3	-18.2	52.2	-18.8	50.5	-19.4	49.2	-19.9	48.0	-20.3
5		g^{II}	82.6		81.0		79.8		78.8		78.0	
	$b_2=35$	g^I	105.0	27.2	103.4	27.6	102.1	28.0	101.1	28.3	100.2	28.5
	$b_2=75$	g^I	69.1	-16.3	67.5	-16.7	66.2	-17.0	65.2	-17.3	64.3	-17.5

Table 5. Examples for Examining the Effect of Demand Variance on Profit

Ex	Number of demand size	Demand size(i)	$p(i)$	$E(i)$	$V(i)$
6	2	1, 5	$p(1) = p(5) = 0.5$	3	4
7	3	1, 3, 5	$p(1) = p(3) = p(5) = 1/3$	3	2.67
8	5	1, 2, 3, 4, 5	$p(1) = p(2) = p(3) = p(4) = p(5) = 0.2$	3	2

Table 6. Comparison of g^I and g^{II} using Examples 1-3 in <Table 5>

Ex			$b_1=10$		$b_1=15$		$b_1=20$		$b_1=25$		$b_1=30$	
				%		%		%		%		%
6		g^{II}	65.2		63.1		61.4		60.0		58.8	
	$b_2=35$	g^I	87.0	33.4	84.8	34.4	83.0	35.2	81.5	35.9	80.3	36.6
	$b_2=75$	g^I	53.1	-18.6	50.9	-19.4	49.1	-20.0	47.6	-20.5	46.4	-21.0
7		g^{II}	65.7		63.6		62.0		60.6		59.4	
	$b_2=35$	g^I	87.6	33.3	85.4	34.2	83.6	35.0	82.2	35.7	81.0	36.3
	$b_2=75$	g^I	53.6	-18.4	51.5	-19.1	49.7	-19.7	48.3	-20.3	47.1	-20.7
8		g^{II}	66.0		63.9		62.3		60.9		59.8	
	$b_2=35$	g^I	87.9	33.2	85.7	34.1	84.0	34.9	82.6	35.5	81.4	36.1
	$b_2=75$	g^I	53.9	-18.3	51.8	-19.0	50.1	-19.6	48.7	-20.1	47.5	-20.6

다음은 LD 크기와 변동성이 모형 I과 II의 수행도에 미치는 영향을 분석한다. 이를 위해, <Table 5>에 수요 크기의 기댓값은 동일하지만 수요 종류와 분산이 상이한 수치 예제들을 구성하였다. 예제 6, 7, 8은 각각 수요 크기가 2가지, 3가지, 5가지 종류로 이루어졌다. 분석의 용이성을 위해 각 예제에 대해 동일한 $p(i)$ 를 갖도록 설정하였다. 예제 6-8에서 $E(i)$ 는 3으로 같지만 $V(i)$ 는 예제 6이 4로 가장 크고, 예제 8이 2로 가장 작다. <Table 6>은 <Table 5>의 3개 예제에 대한 모형 I과 II의 이익 수행도를 보여 준다. <Table 6>에서도 <Table 4>와 동일한 b_1 과 b_2 의 조합을 고려하였다. <Table 6>의 결과를 보면, 모든 b_1 과 b_2 의 조합에 대해 $V(i)$ 가 작아질수록 %값은 줄어드는 하지만 그 감소폭은 크지 않음을 알 수 있다.

6. 결 론

본 논문은 장기공급계약을 이행하고 있는 부품업체가 단기계약주문을 수용하는 문제를 제시하였고, MDP 모형을 사용하여 생산용량할당과 단기계약주문통제를 위한 최적 정책의 구조를 분석하였다. 관련 문헌들과 비교했을 때, 본 논문의 가장 큰 차이점은 중단 불가한 배치생산프로세스와 복합포아송 수요프로세스를 가정한다는 점이다. 중단 가능한 생산에서는 진행 중인 장기계약생산이 중단되고 단기계약생산이 진행되는 상황이 발생할 수 있다. 한 단계 혼성 제조시스템에서는 단기계약주문에 대한 지연재고허용 여부가 중요하기 때문에 지연재고정책의 이익 성과 효과성을 검증하였다. 분석 결과, 단기계약주문매출액 대비 지연재고비용이 일정 수준을 초과하게 되면, 지연재고정책은 기업 이익을 악화시킬 수 있음을 알 수 있었다. 또한, 장기계약수요의 크기와 변동성

이 기업의 이익 성과에 미치는 영향을 분석하였다. 수치 실험 결과, 최소 수요와 최대 수요 발생 확률이 중간 수요들의 발생 확률보다 이익 성과에 미치는 영향이 크다는 사실을 알 수 있었다. 한편, 수요 변동성이 이익 성과에 미치는 영향은 크지 않음을 알 수 있었다. 본 논문의 연구 결과는 혼성 생산 방식을 갖는 제조 기업에게 효과적인 정책 수립을 위한 이론적 토대를 제공할 수 있을 것이다. 추후 연구에서는 최적 생산용량분배와 단기계약 주문수용 정책을 수리적으로 증명하고자 한다.

REFERENCES

- [1] Abedi, A. & Zhu, W. (2020). An advanced order acceptance model for hybrid production strategy, *Journal of Manufacturing Systems*, 55, 82-93.
- [2] Bish, E., Muriel, A., & Biller, S. (2005). Managing flexible capacity in a make-to-order environment, *Management Science*, 51, 167-180.
- [3] Buyukkoc, C., Varaiya, P., & Walrand, J. (1985). The μ -rule revisited. *Advanced Applied Probability*, 17, 237-238.
- [4] Carr, S. & Duenyas, I. (2000). Optimal admission control and sequencing in a Make-to-stock and Make-to-order production system. *Operations Research*, 48, 709-720.
- [5] Cohen, M., Agrawal, N., & Agrawal, V. (2006). Winning in the Aftermarket. *Harvard Business Review*, DOI: 10.1225/R0605H.
- [6] de Vericourt, F. & Karaesmen, Y. (2002). Stock allocation for a capacitated supply chain. *Management Science*, 48, 1486-1501.

- [7] Ding, Q., Kouvelis, P., & Milner, J. (2016). Inventory rationing for multiple class demand under continuous review. *Production and Operations Management*, 25, 1344-1362.
- [8] Duenyas, I. & Van Oyen, M. P. (1996). Heuristic scheduling of parallel heterogeneous queues with set-ups. *Management Science*, 42, 814-829.
- [9] Elhafi, M. (2009). Optimal integrated production and inventory control of an assemble-to-order system with multiple non-unitary demand classes. *European Journal of Operational Research*, 194, 127-142.
- [10] Ha, A. Y. (1997a). Inventory rationing in a make-to-stock production system with several demand classes and lost sales. *Management Science*, 43, 1093-1103.
- [11] Ha, A. Y. (1997b). Stock rationing policy for a make-to-stock production system with two priority classes and backordering. *Naval Research Logistics*, 44, 457-472.
- [12] Hofri, M. and Ross, K. W. (1987). On the optimal control of two queues with server setup times and its analysis. *SIAM Journal of Computing*, 16, 399-420.
- [13] Iravani, S., Liu, T., & Simchi-Levi, D. (2012). Optimal production and admission policies in make-to-stock/make-to-order manufacturing systems. *Production and Operations Management*, 21, 224-235.
- [14] Kim, E. & Van Oyen, M. P. (2021). Joint admission, production sequencing, and production rate control for a two-class make-to-order manufacturing system. *Journal of Manufacturing Systems*, 59, 413-425.
- [15] Kouki, C. & Larsen, C. (2021). Rationing policies in a spare parts inventory system with customers differentiation. *International Journal of Production Research*, 59, 6270-6290.
- [16] Kuthambalayan, T. & Bera, S. (2020). Managing product variety with mixed make-to-stock/make-to-order production strategy and guaranteed delivery time under stochastic demand. *Computers Industrial Engineering*, 147, 106603.
- [17] Liu, Z., Nain, P., & Towsley, D. (1992). On optimal polling policies. *Queueing Systems*, 11, 59-83.
- [18] Peeters, K. & van Ooijen, H. (2020). Hybrid make-to-stock and make-to-order systems: a taxonomic review. *International Journal of Production Research*, 58, 4659-4688.
- [19] Puterman, M. (2005), *Markov Decision Processes*, John Wiley and Sons.
- [20] Wang, Z., Qib, Y., Cuic, H., & Zhang, J. (2019). A hybrid algorithm for order acceptance and scheduling problem in make-to-stock/make-to-order industries. *Computers Industrial Engineering*, 127, 841-852.



김 은 갑

서울대학교 산업공학과 학사/석사

Northwestern University 박사

University of Toronto 경영대학 박사후
과정

삼성SDS 수석컨설턴트

현재: 이화여대 경영대학 교수

관심분야: SCM/유통, Simulation, Stochastic
optimization

SNA를 이용한 제조용 로봇산업 디지털 공급망 가치사슬 분석

김태우* · 서창교**†

*한국로봇산업진흥원, **경북대학교 경영학부

Value Chain Analysis on Industrial Robot Supply Chain Using Social Network Analysis

Tae-Woo Gim* · Chang-Kyo Suh**†

*Korea Institute for Robot Industry Advancement

**School of Business Administration, Kyungpook National University

Industrial robots are manufactured as robot products by utilizing parts including sensors, actuators, controllers and consist of a supply chain provided for end-users through robotics system integrator. This study was analyzed with a social network analysis(SNA) by using corporate trade data to find the characteristics of the industrial robot supply chain in Korea. The network of the industrial robot supply chain consist of 3200 nodes and 3488 edges and the key types of demand industry have found to be automobiles, electrical and electronic equipment, and metal and machinery. The characteristics of the supply chain of industrial robot industry in Korea were analyzed that foreigner-invested robot companies are very influential. The size of sales and importance of the network do not accord with each other. It is analyzed that the companies are classified into the ones related to robots, but their sales in the robot sector is remarkably low or there are transactions with the ones which have nothing to do with them. In order to improve the Korean industrial robot industry's competitiveness, the implementation of the systems for specialized in robot fields and systematic management of information on robot-related companies are required and the existing quantitative policies for supplying and diffusion industrial robots need to be implemented with qualitative policies for utilizing and promoting them.

Keyword : Industrial Robots, Supply Chain, Value Chain, SNA, Network Analysis

† **Corresponding Author** : School of Business Administration, Kyungpook National University, 80 Daehak-ro, Buk-gu, Daegu 41566, Korea, Tel: +82-53-950-5425
Fax: +82-53-950-6247 E-mail: ck@knu.ac.kr

Received: 19 November 2021, **Revised**: 10 December 2021, **Accepted**: 15 December 2021

1. 서 론

우리나라 로봇산업은 2008년 지능형 로봇개발 및 보급 촉진법 제정을 기점으로 다양한 연구개발과 실증 및 보급을 통해 2021년 제조용 로봇밀도 세계 1위, 로봇시장 규모 세계 5위권 국가로 성장하였다(IFR, 2021). 최근에는 AI, 5G 등의 신기술이 로봇에 접목되면서 로봇의 스마트화가 비약적으로 진전되고 활용분야도 급속도로 확대되어 전 세계적으로 로봇산업에 투자, M&A 등이 확대되는 유례없는 로봇 붐이 형성 중에 있다. 로봇산업은 크게 제조용 로봇과 서비스 로봇으로 구분할 수 있으며 제조용 로봇의 경우 정밀기계 산업 기반의 일본과 유럽기업이 세계시장을 주도하고 있는 것으로 분석되고 있다(산업연구원, 2021). IFR(2021)에 따르면 세계 로봇시장은 그간 제조용 로봇을 중심으로 성장해 왔으나 2019년을 기점으로 서비스용 로봇이 제조용 로봇의 시장규모를 역전하였다. 이는 비대면 경제 확대 등으로 서비스용 로봇시장의 성장과 세계적인 경기침체에 따른 투자규모 축소 등의 원인이 있는 것으로 분석되고 있다. 그럼에도 제조용 로봇은 제조현장에서 제품 생산부터 출하까지 공정 내 주요작업을 수행하고 있으므로 중요성이 매우 크며 스마트 팩토리 기술과 융합하여 더욱 발전할 것으로 기대되고 있다.

세계 제조용 로봇시장은 일본과 유럽, 중국 기업이 주도하고 있으며 우리나라 제조용 로봇 기업의 기술적, 경제적 경쟁력은 매우 취약하다(산업통상자원부 & 한국산업기술평가관리원, 2017). 이에 국내 제조용 로봇산업의 경쟁력 향상을 위해 로봇산업의 가치사슬 구성을 분석하여 정책적 지원방안 모색이 필요한 시점이다.

산업은 관련 기업 간의 가치사슬 연계를 통해 공급망을 형성하고 제품기획, 생산, 판매 등의 체계를 중심으로 성장하고 있으며 산업의 종합적인 분석과 발전전략 수립을 위해서는 핵심 기업을 중심으로 한 전후방 공급망 분석이 필요하다. 그간의 산업 공급망 분석에 대한 국내외 사례 및 연구들은 대부분 전문가 평가 또는 설문 등의 정성적인 방법으로 정책적 지원방안을 도출하여 왔으며 데이터를 바탕으로 한 정량적인 방법에 의한 연구는 거의 존재하지 않았다(문영수, 임대현, 2015).

이에 본 연구에서는 국내 제조용 로봇산업 공급망 가치사슬을 제조용 로봇산업 내 기업 간 매출, 매입 연결관계를 사회 네트워크 분석(Social Network Analysis, SNA) 방법을 이용하여 공급망 가치사슬 내 핵심기업을 탐색하고 특징을 분석하여 국내 제조용 로봇 산업의 육성정책

수립을 위한 기초자료 구축에 활용하고자 한다.

본 논문의 구성은 다음과 같다. 제2절에서는 제조용 로봇에 대한 정의와 현황, 공급망 가치사슬의 정의와 분석, 사회 네트워크 분석에 대한 선행연구를 정리하였다. 제3절에서는 연구모형 및 연구주제를 도출하고, 제4절에서는 연구결과를 제시하였으며, 제5절에서는 결론 부분으로 연구결과 요약 및 시사점을 제시하였다.

2. 선행연구

2.1 제조용 로봇의 정의와 현황

제조용 로봇은 기술의 발전에 따라 개념이 일부 상이한 점은 있으나 종합하면 각 산업의 제조현장에서 제품 생산부터 출하까지 공정 내 작업을 수행하기 위한 로봇으로 자동제어 되고 재프로그래밍이 가능하며 다목적인 3축 또는 그 이상의 축을 가진 자동조정 장치로 정의되고 있다(산업통상자원부, 2019).

IFR(2021)에 따르면 2020년 기준 세계 제조용 로봇시장 규모는 약 132억 달러로 연 평균 8% 성장하고 있다(<Table 1> 참고). 중국, 일본, 미국, 독일, 한국 등 제조업 강국이 제조용 로봇 시장의 전체 75%를 차지하고 있으며 핵심수요 산업은 자동차, 전기·전자, 기계·금속으로 분석되고 있다.

Table 1. Global Industrial Robotics Market Size(IFR, 2021)

Category	2020 year (Millions of \$)	CAGR (2015-2020)
Total	24,257	9%
Industrial robots	13,168	8%
Service robots	11,089	10%

산업연구원(2021)은 세계 제조용 로봇시장에서 일본기업이 시장의 절반 이상을 점유하면서 압도적인 경쟁력을 보유하고 있는 것으로 분석하였다. 특히 일본의 3대 로봇기업인 화낙, 가와사키, 야스카와의 2019년 세계시장 점유율은 약 44%에 달한다. 우리나라 로봇산업이 일본과 비교하여 나타나는 차이점은 기술 경쟁력이 있는 기업의 저변이 절대적으로 부족하며 시장규모에 비해 일본과 한국의 기술력과 기업규모의 차이가 현격한 것이다.

제조용 로봇산업의 공급망 가치사슬은 크게 센서, 구동기, 제어기 등의 로봇부품을 활용하여 로봇이 제조되고 로봇 시스템 통합(System Integration, 이하 SI)사업자

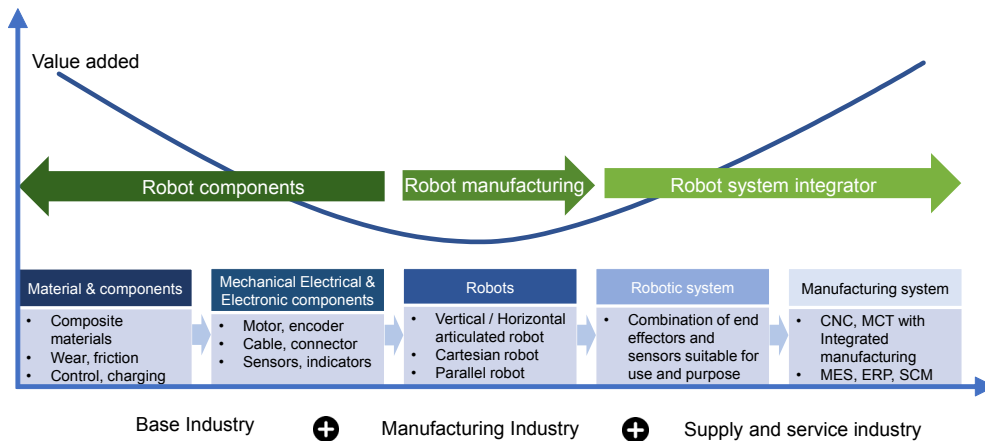


Fig. 1. Supply chain of Industrial robots(로봇사회추진위, 2019)

를 통해 최종 소비자에게 제공되는 형태로 구성된다. 로봇사회추진위(2019)는 로봇 관련 환경변화와 향후 대책 연구에서 <Fig. 1>과 같이 제조용 로봇의 공급망 가치사슬을 로봇부품, 로봇제조, 로봇SI로 나누고 이에 따른 부가가치 분포를 분석하였다. 로봇부품은 타 산업과 공유되는 부품과 로봇전용 부품으로 나눌 수 있으며 정밀구동기, 센서, 제어기 등 일부 부품을 제외하고는 타산업과 공동 사용이 가능하다. 로봇제조 분야는 전통적인 제조용 로봇에서 인간과 협업하는 협동로봇으로 발전하고 있으며 5G, AI기술과의 접목으로 자율제조 환경구현이 가능한 지능화된 제조용 로봇으로 변화하고 있다. 로봇SI는 시스템 설계, 최적의 하드웨어 선정에서 발주·조달, 사용자의 필요에 맞춘 응용 소프트웨어 개발, 시스템의 유지·보수 등 전 과정을 포함하는 것으로 로봇 도입 희망기업의 업무분석을 통해 최적의 로봇시스템을 구축·설치·운전을 도와주는 역할을 수행한다.

한편 EC(2021)는 로봇 가치사슬을 로봇 부품생산, 로봇 생산, 로봇 시스템 통합, 로봇 응용의 단계로 구분하고 EU의 제조용 로봇 가치사슬의 SWOT을 분석하였으며, 산업통상자원부 & 한국산업기술평가관리원(2017)은 제조용 로봇 가치사슬을 다양한 부품의 모듈화를 통해 로봇으로 제조되고 SI사업자를 통해 최종 고객에게 서비스되는 것으로 분석하였다.

2.2 공급망 가치사슬의 정의와 분석

가치사슬(Value Chain, VC)의 정의는 기업 내에서 이루어지는 생산, 물류, 포장, 마케팅, 사후 서비스 등 기업의 전략적 단위 활동을 구분하여 자사의 강·약점을 파악하고 경쟁 기업과 차별화가 가능한 가치 창출 원천을

분석하기 위해 정립한 경쟁우위 평가이론으로 경쟁우위의 기초인 비용 분석을 위해 일련의 기업 내부 활동을 가치사슬로 정의하였다(Poter, 1985). 이후 Shank & Govindarajan(1993)은 기업 간 거래단위인 공급망 VC를 정의하였는데, 공급망 VC는 원료 공급원으로부터 최종 소비자에게 전달되어 제품사용에 이르기까지의 기본적인 가치창출 활동으로 정의하였다. 기업은 가치창출 프로세스인 전체 공급망 VC의 일원이며 산업의 공급망 VC는 기본적인 원료와 부품을 제공하는 공급자들의 가치창출 프로세스에서 출발하는 것으로 다양한 유형의 구매자나 최종 소비자의 가치 창출 프로세스로 이어져 최종적으로 자재의 폐기 및 재활용으로 귀결된다.

Kaplinsky & Morris(2003)은 공급망 VC의 유형을 크게 2가지로 정의하였다. 첫째, 신발, 의류, 가구 및 장난감과 같은 노동 집약적 산업에 많이 나타나는 구매자 주도형(Buyer-driven chains) 공급망 VC로 대규모 소매업체, 마케터, 브랜드 등이 중요한 역할을 하게 되며 제조공장이 다양한 제3국에 위치하고 무역주도 산업화의 패턴을 가지는 것으로 분석하였다. 둘째, 생산자 주도형(Producer-driven chains) 공급망 VC는 다국적 제조기업 또는 자본 및 기술집약적 산업에 나타나는 형태로 자동차, 항공기, 반도체, 중장비 산업이 해당되는 것으로 분석하였다.

공급망 VC 분석은 연구대상이 되는 기업 혹은 산업의 VC 출발점과 VC상의 기업 간 네트워크의 구조를 매핑하고 역할을 분석하는 것이다. 그간 VC에 대한 연구는 주로 전문가 자문 또는 설문방식을 이용하였으나 설문에 의한 방식은 응답자의 조직 내 위치, 지식수준, 조사범위 등에 따라 응답결과가 달라질 수 있어 최근에는 데이터 기반의 VC 분석을 추진하는 사례가 증가하고 있다. 류재호

등(2014)은 우리나라 풍력산업의 가치사슬 구성과 특징에 대해 연구하였으며, 문영수, 임대현(2015)은 기업거래 데이터 기반의 지능형 산업구조분석시스템 구성을 연구하였다. 김장엽 등(2019)은 뿌리산업의 제조혁신 방안에 대한 연구에서 전후방 산업관점에서 기술혁신의 애로사항을 분석하고 핵심 기술로드맵을 제시하였으며, 정재현(2019)은 사회 네트워크 분석을 통한 전자산업 클러스터를 연구하였고, 조진희(2020)는 기업 간 거래 데이터를 바탕으로 충북 지역의 모빌리티산업 생태계를 분석하였다.

산업에 대한 공급망 VC분석은 매우 복잡하고 기업 간 거래는 외부에서 관찰하기 어려운 한계가 있으며 특히, 데이터를 획득하는 것은 현실적으로 더욱 어렵다. Kaplinsky & Morris(2003)는 이론상의 공급망 VC와 현실 세계의 VC 구조를 <Fig. 2>와 같이 나타내며 현실 세계의 VC는 매우 복잡한 구조를 가지고 있다고 강조하였다.

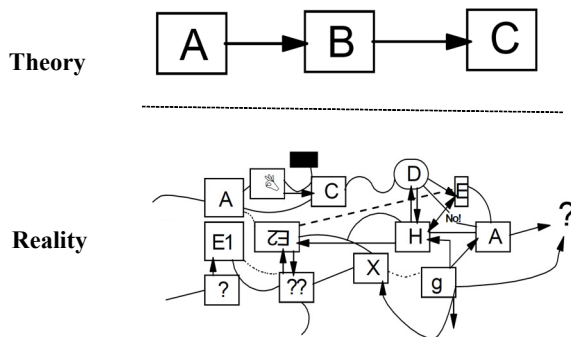


Fig. 2 Value Chain Mapping: Theory and Reality(Kaplinsky & Morris, 2003)

2.3 사회 네트워크 분석

사회 네트워크(Social Network)는 웹(web) 과학의 연구에서 출발한 개념으로 개인 또는 집단이 하나의 노드(node)가 되어 각 노드 간의 상호 의존적인 관계에 의해 만들어지는 사회적 관계 구조를 찾아내기 위한 방법론으로 활용되고 있다. 이는 그래프 이론에 기초하고 있는데 객체 간에 짝을 이루는 관계를 모델링하기 위한 것으로 꼭짓점(node)으로 구성되며 이것들은 관계(edge)를 가지게 된다. 노드는 네트워크 안에 존재하는 개별적인 주체들이며 엣지는 각 노드 사이의 관계를 의미하고, 노드와 엣지 사이의 연결 관계로 네트워크(network)가 형성된다.

사회 네트워크 분석(SNA)은 네트워크 내 행위자들이 맺고 있는 관계를 분석하여 네트워크 구조를 파악하기 위한 정량적 분석 방법으로, SNA 분석을 통해 사회구조 및 주체간의 연결 관계, 집중도 등을 분석하여 수치 및

시각화를 통한 결과 제시가 가능한 연구방법론이다. SNA 분석을 위해 중심성(centrality)을 주로 활용하는데, 중심성은 네트워크 연결망에서 특정 노드가 가지는 영향력을 측정할 수 있어 SNA 분석 지표 중에서 많이 활용되고 있으며, 중심성은 크게 연결중심성, 근접중심성, 매개중심성, 고유벡터 중심성 등으로 구분된다(김용학, 2016).

연결중심성은 한 노드에 연결되어 상호작용하는 엣지의 수로 측정되는 중심성으로 노드의 연결정도를 기준으로 계산되어 각 노드에 연결된 노드의 양이 많을수록 중심성이 높게 나타난다. 근접중심성은 한 노드가 다른 노드들과 평균적으로 얼마나 가까이 있는지를 측정하는 척도로 근접중심성이 높은 노드는 다른 노드들과 전반적으로 가장 짧은 거리에 위치해 적은 단계만으로 여러 노드에 연결될 수 있다. 매개중심성은 두 노드 사이의 최단경로에 존재하는 횟수를 측정하는 것으로 서로 다른 집단 간을 연결하는 노드일수록 높게 나타난다. 매개중심성은 얼마나 많은 노드가 특정 노드를 통과하여 다른 노드와 연결되는 것인지에 관하여 중심성을 부여하여 측정하는 방법으로 네트워크 상에서 노드가 담당하는 매개자 역할의 정도를 나타낸다. 고유벡터 중심성은 측정대상 노드의 중심성과 함께 연결된 다른 노드의 중심성 지표를 고려한 것으로 한 노드의 영향력 또는 중요도를 평가하는데 사용하는 개념으로 네트워크에서 영향력이 높은 중심노드를 나타내고 있다.

SNA 분석방법은 다양한 연구분야에서 활용되고 있는데, Kao(2016)는 기업 간 관계 데이터를 SNA를 활용하여 기업의 공급 네트워크의 생산성 증가, 성과 및 공급망 효율성에 어떤 영향을 미치는지 연구하였다. Zheng et al.(2016)은 건설 프로젝트 관리분야의 연구논문을 대상으로 저자와 연구주제 등에 대해 SNA를 활용하여 연구동향을 분석하였다. Li et al.(2019)은 데이터 분석기반의 기술예측을 위해 SNA를 활용하여 자율주행차 관련 유망 기술을 분석하였으며, Strozzi et al.(2019)은 환자가 의료기관을 선택하는 결정요인 분석을 위해 병원 방문 데이터와 지리 데이터를 활용하여 SNA로 연구하였으며, Kalantari et al.(2021)은 정책연구 보고서와 전문가 검토, SNA를 활용하여 이란의 과학기술 정책결정 네트워크를 분석하였는데 중심성 분석을 통해 각 기관의 역학관계를 측정하였다.

또한 정재현(2019)은 SNA로 국내 전자 업종의 기업 네트워크를 분석하여 97개의 클러스터를 도출하고 매출액과 클러스터 구성 관계를 분석하였다. 정석진 등(2020)은 생태산업단지 구축사업에 참여한 기업의 수요, 공급 DB를 구축하고 SNA를 활용하여 전체 기업 간 네트워크

를 대상으로 특징을 분석하여 자원의 효율적 활용 방안을 모색하였다. 서창교(2020)는 토픽모델링과 SNA를 활용하여 SCM 연구동향에 대해 탐색하였으며 오정진 & 김용진(2020)은 SNA를 활용하여 내연기관차와 전기자동차 부품 공급 네트워크를 비교분석 하였으며, 손용정(2021)은 SNA를 활용하여 광양항을 중심으로 한 정기 컨테이너 서비스의 네트워크 중심성을 분석하여 교역항구를 도출하고 물동량 등 향후 발전방안을 제시하였다.

3. 연구모형

우리나라의 제조용 로봇 활용도는 세계 최상위 수준이나 기술적, 경제적 경쟁력은 매우 취약한 것으로 분석되고 있다. 이미 시장을 석권하고 있는 일본과 유럽 제품은 높은 성능에도 불구하고 대량의 물량을 비교적 낮은 가격에 공급할 수 있는 능력을 보유하고 있으나, 국내 기업들은 기술후발 주자로 낮은 시장 점유율 때문에 부품 수급단가가 현저히 높고 제작단가 또한 높아 가격경쟁력을 확보하기 매우 어려운 상황이다. 이에 국내 제조용 로봇산업의 경쟁력 확보를 위해 산업의 공급망 VC 현황을 분석하여 정책적 지원방안 모색이 필요한 시점이다.

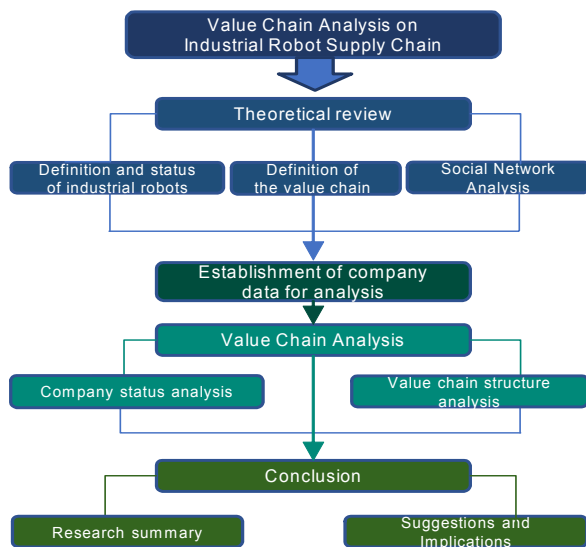


Fig. 3. Research Model

본 연구에서는 <Fig. 3>과 같은 연구모형으로 이론고찰과 제조용 로봇산업 전후방 기업의 일반현황과 매출처, 매입처 DB를 구축하였다. 이를 바탕으로 SNA 중심성 분석방법을 이용하여 공급망 VC 구조를 분석하여 VC 내에는 어떤 기업이 매출, 매입 연결관계가 높은 핵심기

업이며, 이들 핵심기업들 간의 관계와 특징을 분석하여 시사점과 정책제언을 제시하고자 한다.

특정 기업들 간의 거래 규모(매출액 등)에 대한 데이터는 일반적으로 공개되지 않는 현실을 반영하여 본 연구에서는 다른 기업들과의 거래 관계를 활용하여 국내 제조용 로봇산업의 공급망 생태계의 핵심기업을 분석하였다. 기업정보 및 거래처 데이터는 2020년 기준으로 국내에서 로봇관련 사업을 영위하는 기업현황 자료와 매입, 매출의 거래처 데이터를 제공하고 있는 (주)한국기업데이터의 CRETOP DB(<https://cretop.com>)를 기본으로 활용하였으며, 금융감독원 전자공시 시스템(<https://dart.fss.or.kr/>)과 NICE 기업정보(<https://www.nicebizinfo.com>)를 보조적으로 활용하였다. 데이터 분석에는 MS Excel 2016과 Gephi 0.9.2를 활용하였다.

4. 연구방법 및 실증분석

4.1 자료수집 및 표본특성

로봇부품, 로봇제조, 로봇SI로 연결되는 제조용 로봇산업 공급망 VC 단계에 해당되는 기업 발굴을 위해 한국 표준산업분류(KSIC)와 기존 국내에 발표된 정책보고서, 언론보도 등을 참고 하였다.

KSIC 분류상 로봇산업은 ‘산업용 로봇제조업(29280)’의 단일 코드를 사용하고 있다. 이에 로봇제조 업종의 기업은 29280 코드에 해당되는 기업을 CRETOP DB에서 추출하였으며, 로봇부품 기업은 로봇산업 정책보고서, 언론보도를 통해 기업을 발굴하였으며, 로봇SI업종은 2020년 한국로봇산업진흥원에서 발간한 ‘로봇SI 기업편람’에 수록된 기업을 활용하였다. 각각의 세부 기업정보는 CRETOP DB를 활용하였다. 본 연구에 활용한 기업 데이터는 총 570개사이며 공급망 VC분석에는 매출처, 매입처 정보가 제공되는 480개사의 데이터를 분석에 활용하였다(<Table 2> 참조).

Table 2. Data Collection for Analysis

Classification	Data sources	Industry analysis	VC analysis
Robot components	• Media & policy reports	54	49
Robot Manufacturing	• Manufacture of industrial robots(KSIC : C29280)	364	293
Robot System integrator	• Robot SI company handbook	152	138
Total		570	480

매출처, 매입처의 거래 데이터 분석에 있어 동일한 기업명을 가진 기업은 본사 소재지와 대표자명, 사업내용 등을 확인하여 분석대상 기업 포함여부를 판단하였으며, 기업명 표기 방법에 따라 유사기업명은 동일기업 여부를 확인하는 등 데이터 전처리 후 분석에 반영하였다.

4.2 분석대상 기업의 현황

분석대상 기업의 지역별 분포는 수도권(경기/서울/인천)에 317개사(55.6%)가 입지하고 있어 전체 기업의 절반 이상이 수도권에 입지하고 있는 것으로 분석되었다. 세부적으로 로봇부품 기업은 경기>인천>서울>대구 순으로 분석되어 전체 기업의 분포와 유사한 분포로 조사되었고, 로봇제조 기업은 경기>대구>경남>경북 순으로 분석되어 대구/경북의 로봇제조 기업 입지가 수도권을 제외하면 가장 높은 것으로 조사되었다. 로봇SI 기업은 경기>서울>인천>경북 순으로 분석되어 로봇SI 기업 역시 수도권에 집중된 것으로 분석되었다.

조사대상 기업 당 평균 38.2명의 인력이 종사하고 있는 것으로 나타났으며 2020년 기준 국내 제조용 로봇산업의 종사자는 총 19,085명으로 조사되었다. 소기업인 10인 미만이 45.2%(226개사)로 가장 높았으며, 10-49인이 38%(190개사)로 나타났다. 종합적으로는 50인 미만 기업이 83.2%(416개사)로 나타나 국내 제조용 로봇산업은 소기업 중심의 산업으로 분석되었다. 지역별로 수도

권을 제외하면 대구/경북에 약 2,700명이 종사하고 있어 타 지역에 비해 높게 나타났으며, 업종별로는 로봇SI기업의 종사자수가 많은 것으로 분석되었다. 국내 제조용 로봇기업의 67.5%는 일반 법인형태의 기업으로 분석되었으며 개인사업자(13.2%)와 주식회사 중 자산총액이 120억원이 넘어 회계법인으로부터 의무적으로 회계감사를 받아야 하는 외감법인(12.9%)의 순으로 분석되었다. 로봇부품 업종의 외감법인 비율이 로봇제조와 로봇SI 보다 높은 것이 특징적으로 나타났다(<Table 3> 참조).

KSIC 분류로는 로봇부품 기업은 ‘전동기 및 발전기 제조업’에 해당되는 기업이 가장 많았으며, ‘그 외 기타 전자부품 제조업’, ‘기어 및 동력전달장치 제조업’의 순으로 분석되었으며, 로봇제조 기업은 ‘산업용 로봇 제조업’에 가장 많고, 로봇SI 기업은 ‘산업처리공정 제어장비 제조업’, ‘그 외 기타 특수목적용 기계 제조업’, ‘응용 소프트웨어 개발 및 공급업’, ‘반도체 제조용 기계 제조업’의 순으로 나타났다. 부품 제조업의 특성상 로봇전용 부품 제조보다는 다양한 업종에서 활용될 수 있는 부품을 제조하는 기업이 많았고, 로봇SI 기업의 경우 반도체 장비관련 기업이 많은 것이 특징적으로 분석된다.

4.3 공급망 가치사슬 분석

국내 제조용 로봇산업을 구성하는 로봇부품, 로봇제조, 로봇SI 기업의 매입처와 매출처의 데이터를 바탕으로 공급

Table 3. Classification of Industrial Robot Company in Korea

Classification			Location												
			Capital area		Daejeon/Chungcheong		Gwangju/Jeolla		Busan/Ulsan/Gyeongnam		Daegu/Gyeongbuk		Gangwon		
Total			570	317	55.6%	66	11.6%	20	3.5%	84	14.7%	82	14.4%	1	0.2%
Sectors	Robot components		54	38	70.4%	7	13.0%	1	1.9%	2	3.7%	5	9.3%	1	1.9%
	Robot manufacturing		365	185	50.7%	40	11.0%	14	3.8%	66	18.1%	60	16.4%		
	Robot SI		151	94	62.3%	19	12.6%	5	3.3%	16	10.6%	17	11.3%		
Number of workers	less than 10 people		226	115	50.9%	23	10.2%	13	5.8%	35	15.5%	39	17.3%	1	0.4%
	10 to 49 people		190	98	51.6%	30	15.8%	5	2.6%	30	15.8%	27	14.2%		
	50 to 99 people		36	19	52.8%	6	16.7%			5	13.9%	6	16.7%		
	more than 100 people		48	39	81.3%	2	4.2%			3	6.3%	4	8.3%		
	Not mentioned		70	46	65.7%	5	7.1%	2	2.9%	11	15.7%	6	8.6%		
Corporate type	Corporation		383	209	54.6%	50	13.1%	18	4.7%	48	12.5%	57	14.9%	1	0.3%
	Private Company		75	37	49.3%	5	6.7%	2	2.7%	18	24.0%	13	17.3%		
	External audit corporation		73	43	58.9%	7	9.6%			16	21.9%	7	9.6%		
	KOSDAQ		32	24	75.0%	3	9.4%			1	3.1%	4	12.5%		
	KOSPI		3	1	33.3%	1	33.3%					1	33.3%		
	KONEX		1	1	100.0%										
	Not mentioned		3	2	66.7%					1	33.3%				

망 VC를 분석하였는데, 전체 네트워크는 3,200개의 노드와 3,488개의 엣지로 구성되어 있는 것으로 나타났다. VC 내의 핵심기업은 매출, 매입 거래관계에 의한 네트워크 연결정도, 중앙성 분석 등의 SNA 분석방법을 활용하여 네트워크 내 노드(기업)의 중요성을 확인하고자 하였다.

연결중심성(C_D)은 노드의 연결정도를 기준으로 i 는 특정 노드를 나타내며, n 은 네트워크 내에 존재하는 노드의 수를 나타낸다. 특정 노드 i 가 다른 노드 j 와 연결되어 있으면 $a_{ij}=1$ 로 산출되며 연결되어 있지 않은 경우에는 $a_{ij}=0$ 으로 산출되어 이를 표준화하기 위해 $(n-1)$ 로 나누어 준다. 이를 수식으로 표현하면 아래와 같다.

$$C_D(i) = \sum_{j=1}^n a_{ij} / (n-1)$$

전체 네트워크에 대한 연결중심성 분석결과를 시각화 하면 <Fig. 4>와 같다. 연결중심성 분석결과 국내 제조용 로봇 공급망 VC에서 핵심 수요처는 현대·기아 자동차, 삼성전자, 현대중공업, 현대위아 등으로 분석되었

다. 이는 전통적인 제조용 로봇의 핵심 수요업종인 자동차, 전기·전자, 기계·금속 업종에 부합하는 것으로 해석된다.

로봇부품 업종에서는 한국야스카와전기, 삼익THK, 피에조테크놀로지가 연결중심성이 높은 것으로 분석되었으며, 로봇제조 업종에서는 한국화낙, 마이키, 현대로보틱스, 로보스타 등이 높은 연결중심성을 가지고 있으며, 로봇SI 업종에서는 코보시스, 에이비비코리아가 상위 20개 기업에 포함되었다(<Table 4> 참조).

매개중심성(C_B)의 산출은 특정 노드 i 와 네트워크 내에 존재하는 노드의 수(n)를 바탕으로 노드 j, k 가 직접적으로 연결되어 있지 않은 상태에서 특정 노드 i 가 노드 j 와 k 를 연결해 주는 위치에 있게 되면 $g_{jk}(i)=1$, 그렇지 않으면 $g_{jk}(i)=0$ 으로 산출된다. 표준화를 위해 $\frac{2}{(n-1)(n-2)}$ 를 곱해준다. 이를 수식으로 표현하면 아래와 같다.

$$C_B(i) = \left(\sum_{j < k} g_{jk}(i) / g_{jk} \right) \frac{2}{(n-1)(n-2)}$$

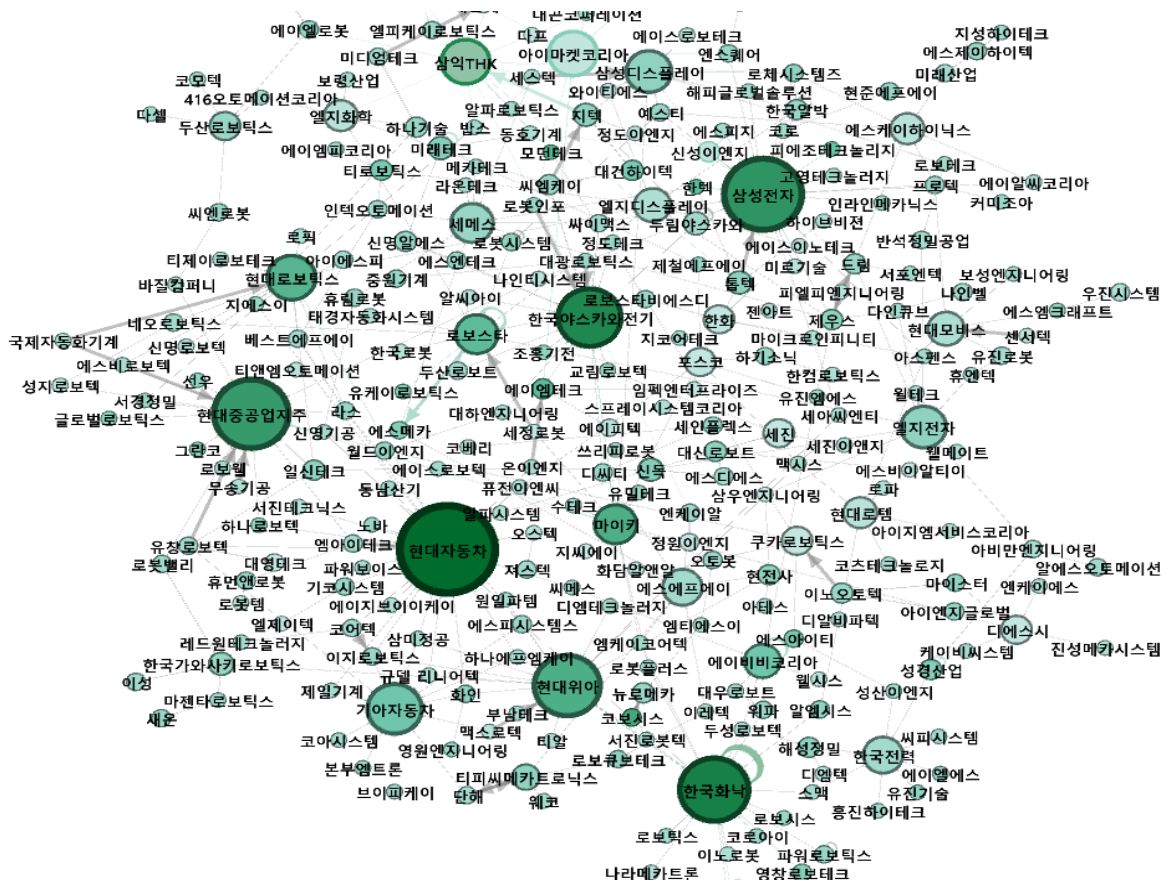


Fig. 4 Result of Degree Centrality

전체 네트워크를 매개중심성 기준으로 분석한 결과, 로봇부품 업종에서는 한국야스카와전기, 삼익THK, 세인플렉스가 매개중심성이 높은 기업으로 분석되었으며, 로봇제조 업종에서는 씨엠케이, 지텍, 마이키, 엔케이알 등의 기업이, 로봇SI 업종에서는 로체시스템즈, 톽텍, 한국알박의 매개중심성이 높은 것으로 분석되었다(〈Table 4〉 참조).

근접중심성은 각 노드 간 거리를 근거로 중심성을 측정하는 지표로서, 한 노드에서 다른 노드에 도달하기까지 필요한 최소 단계의 합으로 정의할 수 있다. 근접중심성 분석결과 각 노드별로 유의미한 차이가 없어 추가로 고유벡터 중심성 분석을 실시하였다. 고유벡터 중심성(λ)은 고유 값(Eigenvalue)과 최적해 산출을 위한 상수인 비례계수 이고, A는 1-mode network, n 은 전체 노드 수를 나타낸다. 고유벡터 중심성 분석을 수식으로 표현하면 아래와 같다.

$$C_E = \frac{1}{\lambda} \sum_{j=1}^n A_{ji} C_i(V_j)$$

고유벡터 중심성을 기준으로 전체 네트워크를 분석한 결과, 〈Table 4〉과 같이 로봇부품 업종에서는 한국야스

카와전기, 로봇제조 업종에서는 두림야스카와, 씨엠케이, 마이키, 한국화낙 등의 기업이 도출되었다. 로봇SI 업종의 기업은 상위 20개사에는 도출되지 않았다. 로봇제조 업종의 두림야스카와는 자회사와의 거래가 많아 중심성이 높게 나타났다. 아울러 수요기업으로는 삼성디스플레이, 엘지화학, 삼성전자 등이 도출되었으며 공급기업으로는 현대위아, 현대중공업지주와 알파엠, 두산인프라코어, 화천기공 등의 공작기계 제조기업이 도출되었다(〈Fig. 5〉 참조).

네트워크 내에서의 커뮤니티 형성 현황을 분석하기 위해 Modularity 분석을 실시하였다. Modularity는 네트워크 내에서 엣지들의 무작위적인 연결들과 비교했을 때 얼마나 더 많은지 정량화 하는 것으로 네트워크 G와 각 노드의 엣지의 개수를 유지한 후 대상을 랜덤하게 변화시킨 네트워크와 원래 네트워크를 비교하여 도출한다. 이를 수식화 하면 아래와 같다.

$$Q = \frac{1}{2M} \sum_{i,j} (a_{ij} - \langle t_{ij} \rangle) \delta[C(i), C(j)]$$

여기서 Q는 모듈성이고 M은 전체 링크 수이며 N은

Table 4. Top 20 Results of Degree, betweenness and Eigenvector Centrality

No	Degree centrality analysis		Betweenness centrality analysis		Eigenvector centrality analysis	
	company names	Degree	company names	Betweenness Centrality	company names	Eigenvector Centrality
1	Hyundai Motor	27	Korea Yaskawa Electric	.000306	Doolim-Yaskawa	1
2	Korea Fanuc	24	Samik THK	.000263	Korea Yaskawa Electric	.918677
3	Korea Yaskawa Electric	23	CMK	.000226	Samsung Display	.716635
4	Samik THK	23	G-Tech	.000182	LG Chem	.657385
5	Samsung Electronics	21	Mykey	.000151	Samsung Electronics	.637088
6	Hyundai Heavy Industries Holdings	20	Rorze systems	.000090	CMK	.604658
7	Cobosys	18	Toptech	.000087	Mykey	.586056
8	Mykey	17	NKR	.000054	Hyundai Wia	.517284
9	Hyundai Wia	17	Daegun Hitech	.000037	Korea Fanuc	.504712
10	Hyundai Robotics	16	Shinsung ENG	.000036	Hyundai Heavy Industries Holdings	.482212
11	Modern tech	16	Shinmyung RS	.000029	LG Display	.456935
12	AM Tech	15	Cymax	.000027	RPM	.450766
13	Robostar	15	Jungwon ENG	.000020	G-Tech	.415401
14	Hantech	15	AM Tech	.000020	Daegun Hitech	.414238
15	Robik	14	Doolim-Yaskawa	.000020	A-ra	.410647
16	Piezo technology	14	Ulvac Korea	.000020	Robostar	.377821
17	ABB Korea	13	Mirae Tech	.000020	Semes	.366494
18	Kia Motors	13	DCT	.000018	Doosan Infracore	.356958
19	Daegun Hitech	12	Robostar	.000011	Hwacheon Machinery	.347450
20	Hanyang Robotics	12	Seinflex	.000010	Gyeongnam Electro-Mechanics	.342696

전체 노드수, a_{ij} 는 i, j 간 링크, t_{ij} 는 각 노드가 가지는 엣지의 개수 등을 나타낸다.

제조용 로봇 공급망 VC 네트워크의 커뮤니티 형성을 알아보기 위한 Modularity 분석 결과, 전체 네트워크에서는 5개 이상의 노드와 관계를 맺고 있는 커뮤니티가 총 77개 형성되어 있는 것으로 분석되었다. Degree가 높은 노드를 중심으로 관련 기업과 커뮤니티를 형성하고 있으며 Degree가 높은 복수의 기업이 참여하는 커뮤니티는 다수의 서브 커뮤니티와 상호 연결되어 매우 복잡한 관계를 형성하고 있는 것으로 분석되었다. 그러나, 소규모 커뮤니티 내의 기업 간에는 연결성이 강하고 다른 커뮤니티에 속한 기업과는 상대적으로 연결성이 약해 모듈성 지수에 따른 세부 네트워크 특징을 분석하기에는 한계가 있다.

제조용 로봇 공급망 VC에 대한 전체 네트워크 중심성 분석결과를 종합하면 첫째, 로봇부품 분야에서는 일본 야스카와전기가 100% 지분을 보유한 한국야스카와전기의 중요도가 매우 높은 것으로 분석되었으며, 일본THK

에서 약 33%의 지분을 보유하고 있는 삼익THK도 핵심기업으로 분석되었다. 그 외 액츄에이터 전문기업인 피에조테크놀로지, 엔코더 전문기업인 세인플렉스 등도 핵심기업으로 도출되었다.

둘째, 로봇제조 분야에서도 일본 화낙이 약 94%의 지분을 보유한 한국화낙과 일본 야스카와전기에서 약 35%의 지분을 보유한 두립야스카와의 중요도가 높은 것으로 분석되었다. 또한 일본 야스카와 로봇의 국내 판매기업인 마이키, 한국야스카와전기-지텍 등과 네트워크 관계에 있는 씨엠케이, 두립야스카와 및 삼익THK과의 네트워크 관계에 있는 지텍, 지텍·씨엠케이 등과 네트워크 관계에 있는 대건하이텍이 중요기업으로 분석되었다. LG 전자 계열사이고 코스닥 상장기업인 로보스타도 중요기업으로 분석되었다.

셋째, 로봇SI 분야에서는 스위스 에이비비(ABB)에서 100% 지분을 보유한 에이비비코리아, 일본 로체에서 약 40%의 지분을 보유한 로체시스템즈가 중요기업으로 분석되었으며, 자동화 전문기업인 톱텍과 코보시스도 중요

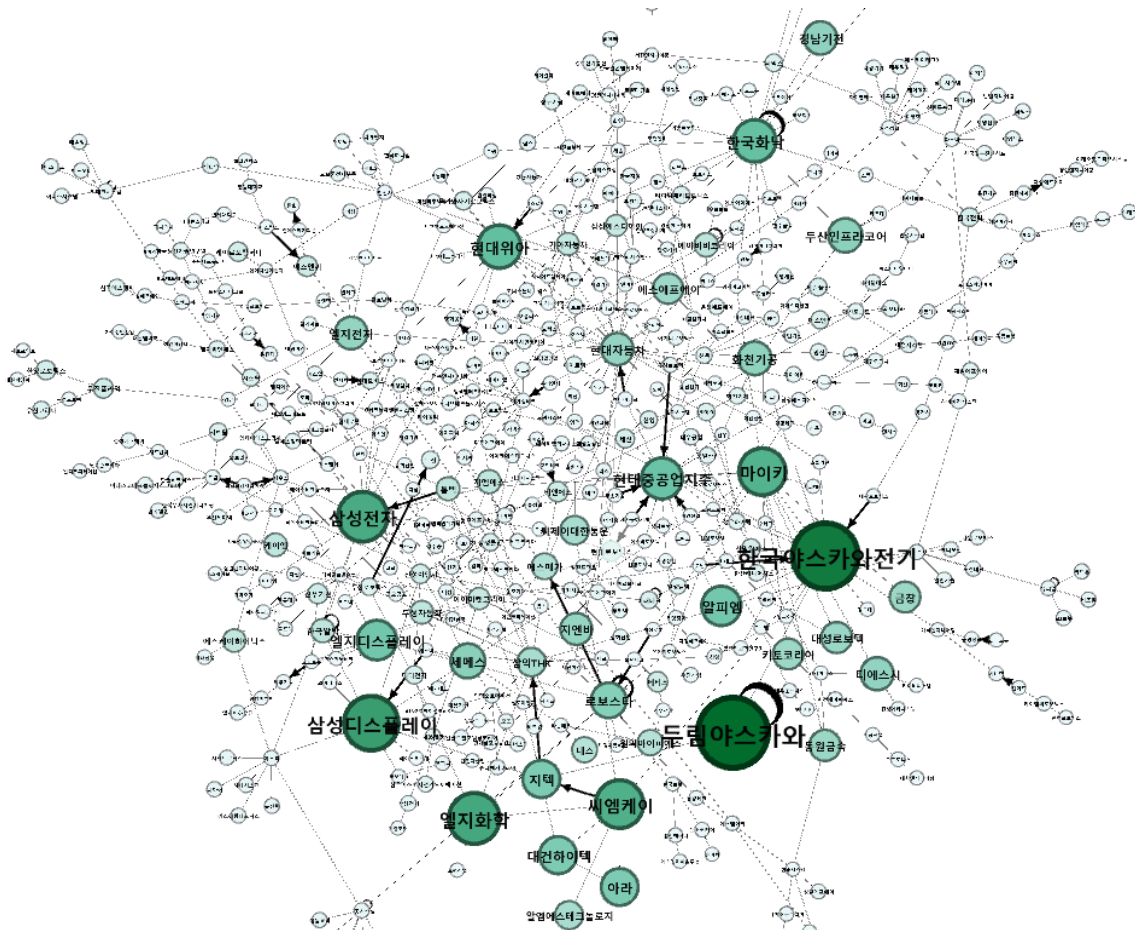


Fig. 5. Result of Eigenvector centrality

기업으로 분석되었다.

전체 네트워크 중심성 분석내용을 종합하면 국내 제조용 로봇 공급망 VC는 일본 관련기업의 기술력에 많이 의존하고 있는 것으로 해석된다. 일본 기업에서 국내 투자한 기업, 그 기업과 거래 관계를 맺고 있는 기업들이 네트워크에서 다양한 기업과 거래관계를 맺고 있는 것으로 분석되어 전체 네트워크에서 일본 관계기업의 영향력이 높은 것으로 분석된다.

4.4 세부 업종별 공급망 가치사슬 특징

전체 네트워크에서 로봇부품, 로봇제조, 로봇SI 업종을 추출하여 세부 업종별 공급망 VC 구성의 특징을 살펴보았다.

첫째, 로봇부품 네트워크는 총 376개의 노드와 355개의 엣지로 구성되어 있는데, 전체 네트워크 분석에서 도출된 로봇부품 분야 핵심기업인 피에조테크놀로지, 삼익THK, 세인플렉스와 함께 로보큐브테크, 서보산전, 효성감속기, 엔티렉스 등이 핵심기업으로 도출되었다(Fig. 6) 참조). 로봇부품 기업들의 중심성은 대부분 Out-degree가 높게

나타나 공급망 상위 기업으로 공급하는 부품업종의 특징이 반영된 결과로 해석된다. 로봇부품 업종 내 네트워크에서 커뮤니티는 26개로 관련 기업과의 거래를 유지하면서 소재·부품단위에서 모듈형태로 제품화 시키는 것으로 분석되었다. 또한 로봇부품 업종의 중심성 상위 20개사를 대상으로 기업부설연구소 운영현황을 살펴본 결과, 모두 부설연구소를 보유하고 있는 것으로 분석되었다. 이는 오랜 시간 연구개발을 통해 제품화가 가능한 부품소재의 특징이 반영된 것으로 분석된다. 다만, 한국야스카와전기와 같이 네트워크 중심성이 높은 외국인 투자 기업은 기업부설연구소를 운영하지 않고 이노비즈 등과 같은 기업인증 없이 주로 기술지원과 마케팅 업무를 중심으로 수행하는 것으로 분석된다.

둘째, 로봇제조 네트워크는 총 2015개의 노드와 2127개의 엣지로 구성되어 있는데, 수요기업인 현대자동차, 현대중공업지주와 한국화낙, 모던테크, 로보스타, 한택, 에이엠테크, 현대로보틱스, 로빅 등의 로봇제조 기업, 한국야스카와전기 등의 로봇부품 기업이 핵심기업으로 도출되었다. 수요기업의 중심성은 In-degree로만 구성되어 있고 로봇제조 기업은 In-degree와 Out-degree가 함

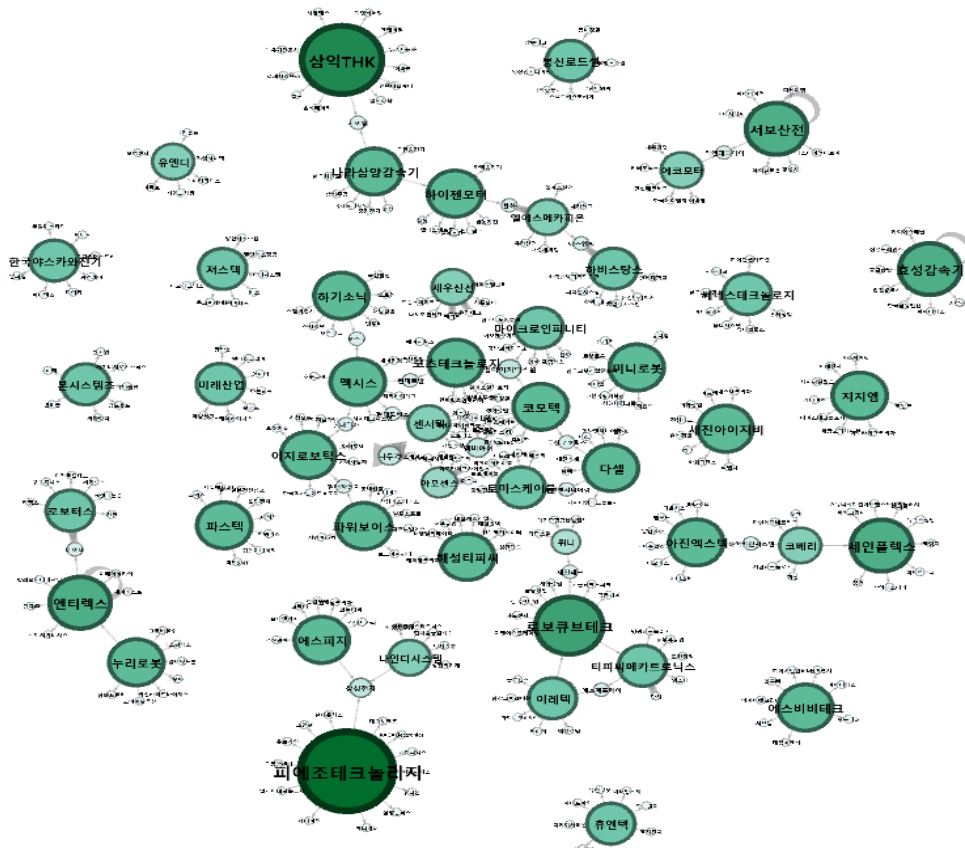
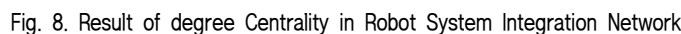


Fig. 6. Result of Degree Centrality in Robot Component Network

셋째, 로봇SI 네트워크는 총 1027개의 노드와 1005개



의 엣지로 구성되어 있는데, 연결중심성 분석에서는 삼성전자, 에스아이티, 건우전장, 신성이엔지, 동양기술 등의 기업이 도출되었으며, 매개중심성에서는 제우스, 에이비비코리아, 화인, 프로텍, 알에스오토메이션 등이, 고유벡터 중심성 분석에서는 에이비비코리아 외에 삼성전자, 삼성디스플레이, 엘지디스플레이의 연결성이 높은 것으로 나타났는데 이는 로봇SI 기업에서 전자업종의 주요산업에 로봇공급이 많았고 로봇SI 기업이 반도체 관련 장비와 제조공정 관련 업무를 함께 수행하는 것에 따른 것으로 해석된다(Fig. 8) 참조).

4.5 매출액과 네트워크 분석 비교

분석대상 기업의 매출액과 공급망 VC내의 중요도를 비교·분석하였다. 분석결과 분야별 매출액 상위 기업은 네트워크 상에서 큰 영향력은 없는 것으로 분석되었다. 첫째, 로봇부품 분야의 매출액 1위 기업인 (주)나무가의 경우 3D 카메라 모듈이 자율주행, 로봇 등에 공급이 가능할 것으로 분석되어 로봇부품 기업으로 분류소개 되고 있으나, 매출의 99%('20년 매출기준)를 삼성전자 스마트폰 및 노트북용 카메라 모듈이 차지하고 있어 제조용 로봇 공급망 VC에서는 영향력이 낮은 것으로 분석되었다.

둘째, 로봇제조 분야 매출액 1위 기업인 한국화낙(주)는 공급망 VC내에서의 중요도가 높은 것으로 분석되었으나, 매출액 2위 기업인 한화정밀기계는 반도체 장비 및 공작기계를 주력으로 하고 로봇분야는 (주)한화/기계에서 수행하고 있어 네트워크 상에서의 중요도가 낮은 것으로 분석되었다. 또한 매출액은 높으나 자회사와의 거래로 인한 부분과 거래데이터가 제공되지 않는 경우도 분석결과에 영향을 미친 것으로 해석된다. 전체적으로는 로봇제조 분야 매출액 상위 기업은 로봇제조 네트워크 내에서 전반적인 영향력이 높은 것으로 나타났다.

셋째, 로봇SI 분야는 기업편람 등을 통해 로봇SI 전문 기업으로 분류되고 있으나, 기계가공, 의료 등이 주력분야로 로봇 외 타 분야에서의 주력 매출이 발생하여 네트워크 분석에서는 중요도가 낮은 것으로 분석된다.

이는 로봇기업으로 분류되고 있는 많은 기업들이 로봇 외 분야를 주력으로 하면서 로봇 관련 제품 또는 서비스를 일부 제공하고 있는 것으로 나타나 체계적인 기업 정보의 관리가 필요한 것으로 분석된다. 현재 국내 로봇산업에 대한 분석은 매년 로봇산업 실태조사가 시행되고 있으며 이를 통해 국가 전체 로봇산업의 거시적인 변화와 현상의 분석에는 효과적이나, 실태조사는 모수추정

방식의 조사로 세부적인 기업정보 분석이 어려워 세밀한 정책개발에는 한계가 있는 것으로 분석된다.

5. 결 론

5.1 연구결과 요약

본 연구는 국내 제조용 로봇산업의 공급망 가치사슬을 사회 네트워크 분석방법으로 분석하여 가치사슬 내 핵심기업과 특징을 분석하였다.

첫째, 표준산업분류상 국내 제조용 로봇산업의 로봇부품 기업은 '전동기 및 발전기 제조업(C28111)'에 해당되는 기업이 가장 많았으며, 로봇제조 기업은 '산업용 로봇 제조업(C29280)', 로봇SI 기업은 '산업처리공정 제어장비 제조업(C27216)'에 해당되는 기업이 가장 많았다. 다만, 부품제조업의 특성상 로봇전용 부품 제조 보다는 다양한 업종에서 활용될 수 있는 부품을 제조하는 기업이 많았고 로봇제조와 로봇SI 분야의 기업은 자동화 설비, 반도체 장비, 기계 관련 기업이 많은 것이 특징으로 분석되었다.

둘째, 국내 제조용 로봇산업의 공급망 가치사슬은 생산자 주도형 가치사슬로 자본 및 기술집약적이며 다국적 제조기업의 시장 지배력이 강한 것으로 분석된다. 아울러 제조용 로봇의 핵심 수요처는 현대기아자동차, 삼성전자, 현대중공업, 현대위아 등으로 분석되어 전통적인 제조용 로봇산업의 주요산업으로 분류되는 자동차, 전기·전자, 금속기계 산업에 부합되는 결과로 도출되었다.

셋째, 국내 제조용 로봇 공급망 가치사슬에서 국내 기업의 네트워크 영향력은 낮고 외국인 투자 로봇기업의 네트워크 영향력이 높은 것으로 분석되었다. 로봇부품 분야에서는 한국야스카와전기, 삼익THK 등, 로봇제조 분야에서는 한국화낙, 두립야스카와 등, 로봇SI 분야에서는 에이비비코리아, 로체시스템 등의 기업이 영향력이 높은 것으로 분석된다. 매출액 상위 기업의 네트워크 중심성을 비교 분석한 결과, 매출액 규모와 네트워크 중요도는 일치 하지 않는 것으로 분석되었는데, 이는 로봇관련 기업으로 분류되고 있으나 해당 기업의 로봇분야 매출액이 현저히 작거나 거래 데이터가 로봇 외 분야 기업과의 데이터가 제시 되어 발생한 것으로 분석된다.

특히, 국내 제조용 로봇산업의 공급망 가치사슬 분석 결과의 가장 큰 특징은 일본 관련 로봇기업의 가치사슬상의 중요도가 높은 것으로 분석되었다. 이는 우리나라 제조용 로봇산업의 태동과도 관련이 높을 것으로 해석할

수 있는데, 고병준(1982)의 연구에 따르면 국내에서 가장 먼저 외국산 제조용 로봇을 도입하여 사용한 곳은 현대 자동차 울산공장으로 1978년 일본의 토요타로부터 다점 용접 로봇을 도입하여 자동차 공정 컨베이어에 일렬로 설치하여 운용하였다. 국내 제조용 로봇제조는 1980년 초 일본기업과의 기술제휴를 통해 도입한 것이 출발점이 되고 있는데, 이후 다양한 일본기업과 국내기업과의 기술제휴를 통해 국내 로봇 및 자동차 시장에 진출하였다.

일본 로봇기업의 기술력이 국내 로봇기업과 차이가 현격하고 시장점유율 또한 높은 것으로 분석하고 있는 산업연구원(2021), 산업통상자원부 & 한국산업기술평가관리원(2017)등 선행연구 결과를 고려하면 이러한 기술력 차이가 국내 제조용 로봇산업의 공급망 가치사슬 구성에도 영향을 미친 것으로 해석된다.

5.2 연구 시사점

최근 중국은 자국 제조용 로봇시장 내수의 70%가 외국계 기업에서 생산되는 제조용 로봇인 것에 대한 문제점을 인식하고 자국 제조용 로봇산업 경쟁력 강화를 위한 계획을 추진 중에 있다. 우리나라도 중국의 이러한 움직임 못지 않게 국내 로봇산업의 육성을 위해 다양한 정책적 수단을 이용하여 국내 제조용 로봇기업의 경쟁력 확보를 위해 노력 중에 있다. 그러나 후발주자로 시장에서 경쟁력을 확보하기 위해서는 종합적이고 전략적인 정책적 접근이 필요할 것으로 분석된다.

본 연구를 통한 학문적 시사점으로는 첫째, 로봇산업의 체계적인 분석을 위한 산업분류의 세분화가 필요하다. 현재 KSIC 코드분류에서 로봇산업은 단일 코드인 '산업용 로봇 제조업'(C29280)에 해당하고 있다. 국내 로봇산업의 규모가 작고 관련 전문기업이 많지 않아 KSIC 코드분류의 세분화가 단기간에는 어렵겠으나, 장기적으로는 맞춤형 정책개발을 위해서 자동차 산업과 같이 엔진, 차체, 신품, 재제조 부품으로 세분화하여 체계적인 산업분석이 가능한 토대를 마련할 수 있는 산업범위 세분화가 필요하다.

둘째, 국내 제조업 로봇산업의 기술경쟁력 향상을 위한 지원정책이 필요하다. 국내 제조업 로봇 공급망 가치사슬은 외국인 투자 기업의 영향력이 높은 것으로 분석되었는데 외국인 투자 뿐 아니라 외국인 투자기업과 관계를 형성하고 있는 기업은 연구소 운영 등 기술혁신 활동이 상대적으로 부족한 것으로 나타났다. 특히, 로봇 제조 분야의 연구소 운영 비율이 낮은 것으로 분석되는

데 장기적으로 국내 제조용 로봇산업의 기술경쟁력 향상을 위해서는 기업부설연구소의 설치·운영을 지원할 수 있는 지원정책이 요구된다.

셋째, 규모의 경제 실현이 가능한 로봇전문 기업의 육성이 필요하다. 본 연구는 기업 간 매입처, 매출처 데이터를 바탕으로 정량적 연구를 진행하였는데 국내 로봇 관련 기업은 로봇 제품만을 핵심 매출로 하는 기업이 많지 않아 매우 복잡한 구조의 공급망 가치사슬이 형성되어 있는 것으로 나타났다. 일본, 유럽 기업이 이미 시장을 선점하고 있어 로봇제품 만으로는 기업경영 활동에 한계가 있어 유사분야의 업종을 영위하면서 로봇 관련 제품을 생산 또는 서비스 하고 있는 것으로 분석된다. 국내 로봇산업의 활성화와 글로벌 경쟁력을 높이기 위해서는 지능형 로봇개발 및 보급 촉진법에서 명시하고 있는 로봇전문기업 지정제도의 시행이 필요하다. 로봇산업에 대한 투자는 대규모 비용이 소요되고 대기업도 로봇에 대한 투자를 증대시키고 있다는 점을 고려할 때 로봇기업의 특성화와 규모 확대를 위한 제도의 시행이 필요한 시점이다.

제조용 로봇의 시장에서 국내 기업들이 경쟁력을 확보하기에는 매우 어려운 상황이다. 본 연구를 통한 실무적 시사점으로는 첫째, 우리나라의 강점으로 평가 받는 ICT 융합기술을 활용하여 수요업종별 맞춤형 솔루션 개발을 통한 특화분야 육성이 필요하다. 아주 단순한 작업은 전용장비가 유리하고 복잡한 공정은 로봇기술의 한계로 작업자에 의한 생산이 효과적이다. 로봇은 중간 정도의 복잡한 공정과 제품생산에 적합하므로 수요업종별 특화 솔루션 개발이 매우 중요하다. 이에 5G, IoT, AI, CPS(Cyber Physical System) 등의 기술을 기반으로 인간과 협업이 가능한 협동로봇, 자율공장과 같은 고도화된 유연생산기술 분야의 경쟁력 확보가 가능할 것으로 전망된다.

둘째, 기존 공정장비와 로봇의 협업모델 개발을 통한 로봇산업 경쟁력 확보가 필요하다. 주조 및 용접, 조립 등 부품제조 현장에서는 자동화가 꾸준히 진행되고 있으나, 다품종 소량생산이 많은 기계가공 분야에서는 여전히 수작업에 의존하는 작업이 많이 남아 있다. 이에 국내 공장기계 제조사와 제조용 로봇 제조사와의 협업을 통해 국내 기계가공 제조현장에 최적화된 로봇협업형 공작기계 개발을 촉진하고 AGV(Automated Guided Vehicles)와 AMR(Autonomous Mobile Robot) 등을 공정 내 물류시스템과 연계하여 한국형 스마트 팩토리 구현을 앞당길 필요가 있다.

본 연구는 국내 제조용 로봇산업의 공급망 가치사슬 분석을 위해 로봇제조 업종을 중심으로 전후방 업종을 함께 분석하여 핵심기업과 특징을 분석하였으나, 로봇 기업 중에서 거래데이터가 제공되지 않은 기업은 분석에 포함되지 못한 점이 본 연구의 한계점이라고 할 수 있다. 향후 기업정보의 체계적인 관리시스템 구축으로 로봇기업 데이터가 추가로 확보된다면 보다 심도 있는 연구가 가능할 것으로 기대된다.

REFERENCES

- [1] 고병준(1982). 산업용 로봇의 역사와 현재의 동향. 전자공학회잡지, 9(4), 5-8.
- [2] 김용학(2016). 사회 연결망 분석. 박영사
- [3] 김장엽, 한재현, 정석재(2019). 뿌리산업 중소기업체의 제조혁신 방안 및 전후방산업으로의 파급효과. 한국 SCM학회지, 19(2), 59-73.
- [4] 로봇사회추진위(2019). 로봇による 社会変革 推進会議. 로봇を取り巻く環境変化と今後の施策の方向性. (https://www.meti.go.jp/shingikai/mono_info_service/robot_shakaihenkaku/index.html)
- [5] 류재호, 최타관, 박중구 (2014). 한국 풍력산업의 가치사슬 및 가치시스템 분석. 에너지공학, 23(1), 46-57.
- [6] 문영수, 임대현 (2015). 기업거래데이터 기반의 지능형 산업구조 분석 시스템. 한국기술혁신학회 학술대회 논문집, 368-374.
- [7] 산업연구원 (2021). 미래전략산업 브리프 제18호 : 주요 신산업의 2020년 세계시장 점유율.
- [8] 산업통상자원부 (2019). 제3차 지능형 로봇 기본계획.
- [9] 산업통상자원부, 한국산업기술평가관리원 (2017). 대한민국 로봇산업 기술로드맵.
- [10] 서창교 (2020). 토포모델링과 소셜 네트워크 분석을 이용한 공급사슬관리 연구 탐색. 한국SCM학회지, 20(1), 36-51.
- [11] 손용정 (2021). SNA를 이용한 정기 컨테이너 서비스 네트워크 분석 - 광양항을 중심으로. e-비즈니스 연구, 22(1), 179-194.
- [12] 오정진, 김용진 (2020). 사회연결망 분석을 이용한 내연기관차와 전기자동차 부품 공급 네트워크 비교 분석. 한국SCM학회지, 20(3), 1-13.
- [13] 조진희 (2020). 네트워크 분석을 통한 충북 모빌리티 산업 생태계 분석. 충북연구원.
- [14] 정석진, 정석, 전호석, 임춘성 (2020). 사회네트워크 분석기법을 활용한 생태산업단지 구축사업 네트워크 특성 분석. 한국자원공학회지, 57(2), 168-175.
- [15] 정재현 (2019). 사회적 네트워크 분석을 이용한 전자산업 클러스터 연구. 한국콘텐츠학회논문지, 19(5), 48-63.
- [16] 한국로봇산업진흥원, 한국로봇산업협회 (2020). 국내 로봇SI기업 편람.
- [17] European Commission (2021). *Robotics for Food Processing and Preparation: Advanced Technologies for Industry - Product Watch*.
- [18] Institute of Management Accountant (1996). *Value Chain Analysis for Assessing Competitive Advantage*.
- [19] International Federation of Robotics (2021). *World Robotics 2021 - Industrial Robots*.
- [20] Kalantari, E., Montazer, G., & Ghazinoory, S. (2021). Mapping of a science and technology policy network based on social network analysis, *Journal of Entrepreneurship, Management and Innovation*, 17(3), 115-147.
- [21] Kao, T. W. (2016). *An Empirical Study of Supply Network: Combining Social Network Analysis with Economic Approaches*, Ph.D. Dissertation, State University of New York.
- [22] Kaplinsky, R. and Morris, M. (2003). *A Handbook for Value Chain Research*, IDRC.
- [23] Li, S., Garces, E. and Daim, T. (2019). Technology forecasting by analogy-based on social network analysis: the case of autonomous vehicles, *Technological Forecasting and Social Change*, 148, 119731.
- [24] Porter, M. E. (1985). *Competitive Advantage: Creating and Sustaining Superior Performance*, The Free Press.
- [25] Shank, J. K. & Govindarajan, V. (1993). *Strategic Cost Management - The New Tool for Competitive Advantage*, The Free Press
- [26] Strozzi, F., Garagiola, E., & Trucco, P. (2019). Analysing the attractiveness, availability and accessibility of healthcare providers via social network analysis. *Decision Support Systems*, 120, 25-37.
- [27] Zheng, X., Le, Y., Chan, A.P.C., Hu, Y., & Li, Y. (2016). Review of the application of social network analysis in construction project management research. *International Journal of Project Management*, 34(7), 1214-1225.



김 태 우

경북대학교 경영학박사

현재: 한국로봇산업진흥원

산업정책팀 책임연구원

관심분야 : SCM, 정보시스템 개발 및
운영, 성과분석



서 창 교

경북대학교 경영학사

POSTECH 공학석사

POSTECH 공학박사

현재: 경북대학교 경영학부 교수

관심분야: 지능정보시스템, SCM

재생에너지 생산시스템에서 검사 자동화의 효과 연구*

Mitali Sarkar** · 서용원***†

중앙대학교 ITRC 연구센터, *중앙대학교 경영경제대학 경영학부

Significance of Autonomated Inspection for a Renewable Energy Production System*

Mitali Sarkar**, Yong Won Seo***†

**Information Technology Research Center, Chung-Ang University,

***Department of Business Administration, College of Business and Economics, Chung-Ang University

With the changing lifestyle of humans, the use of technologies increases, thus increasing the necessity of energy demand. Owing to limited nonrenewable energy sources, the demand for renewable energy is growing gradually. On the other hand, the waste generated increases with the increasing population rate. Wastes can be a useful source of energy demand, and environmental pollution can be reduced. This study explains a renewable energy production system. The proposed production system uses wastes as a source of raw material. Wastes are collected and transported from the collection center of a city to produce renewable energy. An autonomated inspection policy is used in the production system. The random rate of unusable wastes is scraped from the system, and the rest are sent for energy production. The total cost of the energy conversion plant is optimized analytically by a classical optimization method. Three numerical examples are given for uniform, triangular, and beta distribution of random scrap rate. The results show that adopting beta distribution for the proposed study is required to obtain the maximum amount of renewable energy at the minimum cost. The transportation cost is found as the most sensitive cost parameter among all cost parameters. The proposed study fulfills both the aims of renewable energy production and waste management.

Keyword : Supply chain, Distributive justice, Behavioral analysis

* This research was supported by the MSIT (Ministry of Science and ICT), Korea, under the ITRC (Information Technology Research Center) support program(IITP-2021-2020-0-01655) supervised by the IITP (Institute of Information & Communications Technology Planning & Evaluation).

† **Corresponding Author** : Department of Business Administration, College of Business and Economics, Chung-Ang University, Seoul 06974, South Korea.
Tel: +82-2-820-5580, E-mail: seoyw@cau.ac.kr

Received: 24 November 2021, **Accepted**: 13 December 2021

1. Introduction

Due to the increasing utilization rate of smart technologies, the energy demand is rising. But there is a limited source of nonrenewable energy. Thus, renewable energy is another solution to this problem. There are different renewable energy sources like solar, wind, water, hydrogen, and waste. Among those sources of renewable energies, wastes are the most convenient because they have very little purchasing or collection cost as raw material and are easy to get. With the rising population, wastes generation is growing up continuously. Those wastes can be used to produce renewable energy, which can be facilitated in two ways: the first one is that the energy demand can be fulfilled properly. The second one is that wastes are managed properly, and pollution can be reduced. Sarkar et al. (2021) established a supply chain network to produce bioenergy from agricultural wastes.

There are two types of manufacturing systems available to produce renewable energy. Traditional manufacturing can provide renewable energy with less cost, but there is a limitation of the capacity of the production system. Smart production can provide renewable energy without having the interruption during production duration. But the investment in smart production is rather high than the traditional production system. The population of a city is known. Thus, the demand pattern of energy consumption is also known. Ullah et al. (2021) recommended with the traditional production system, provided the energy consumption demand is known. The demand may vary with time. Saha et al. (2021) showed that the dynamical investment may need to control the excess energy consumption demand. Rasti-Barzoki and Moon (2021) described a case study on competition between electrical vehicles (EV) and gasoline-based vehicles. They explained the subsidy and incentive policy, which the government will pay due to EV use.

Thus, the research gap exists in the following direction: producing renewable energy to fulfill the market's demand without shortages, and major investments such as the small manufacturing industries (SME) can participate in this production system. A traditional production system can be used instead of a smart production system to avoid more expenses of the production process. From a long-time data source, a city's energy demand can be known and is considered constant when there is no big change. This energy demand can be ful-

filled by the renewable energy produced from the city's wastes. Vandana et al. (2021) established energy effectiveness in a traditional production system. They proved that renewable energy is much more beneficial than conventional energy. Leeuwen et al. (2021) worked with urban biowaste to produce electricity as renewable energy.

Bhuniya et al. (2021) proved that the maintenance of the smart production might need a huge amount of energy consumption. Thus, they recommended a traditional production system rather than the smart production system, provided there is a huge labor cost. If the labor cost is too high, smart production is recommended for the renewable energy production system. Habib et al. (2021) explained the bio-diesel production system as another renewable energy source to maintain sustainability. They considered a supply chain network and different zones to collect their raw materials. Sarkar (2019) developed a multi-stage production system with a traditional production system with a random defective rate of products. As a result, there is a probability that the production may be greater or less than the constant demand. But Sarkar (2019) did not consider any specific product type. That was a general traditional production model.

That model was extended by Sarkar and Sarkar (2020), who considered a multi-stage production system for a bio-fuel production system as another renewable energy resource. But they utilized a smart production system with some major investment but with a reduced setup cost. The set cost reduction was meaningful as the smart production needs a huge setup cost. Thus, if a traditional production system is utilized, the investment may not be required. If the automated inspection is used, then getting defective products is also reduced. Thus, the recommendation is moved towards a traditional production system if the SME does not want to invest huge funds for making a smart production system for automated operation. Ramos and Rouba (2020) investigated a study on renewable energy from solid waste. They proved that renewable energy is the alternative to fossil fuel. Rasti-Barzoki and Moon (2021) considered a supply chain model with a car manufacturer and energy supplier under a governmental policy. They optimize energy consumption and pollution minimization.

Ahmed and Sarkar (2019) introduced renewable energy as the next-generation energy under a basic supply network. The major concern of their research is what type of energy can solve the issue of energy consumption's demand. They

recommended that bio-energy is the best choice among all energy production systems. But there was no specification about which bio-energy can be utilized for the future generation. Also, there is no mention of the perfect raw material for future energy resources. Danish et al. (2019) explained the energy production from municipal solid wastes. But, they did not discuss the inspection policy. Safarzadeh and Rasti-Barzoki (2019a) established a supply chain model for energy-efficient appliances and discussed consumer behavior and government policies. Safarzadeh and Rasti-Barzoki (2019b) analyzed a duopolistic supply chain for energy production. They explained the pricing, government policies, and rebound effects of industry. Lee and Moon (2014) designed a wireless sensor that is energy efficient. They found the optimum energy-efficient path to minimize energy consumption. Ducharme (2010) analyzed the facility of plasma-assisted waste to energy process.

The organization of this study is Section 2 contains the problem description and assumptions of the study. Section 3 explains the model formulation and solution methodology of this research. Numerical experiments and sensitivity analysis are described in Section 4. The last, Section 5, concludes this study.

2. Problem Description and Assumptions

The proposed study aims to produce renewable energy from waste to reduce pollution and save the environment in environmental sustainability. A renewable production system is considered to accomplish the energy demand (d_e). The system works with wastes and generates renewable energy with a constant production rate (P_e). The manufacturer collects wastes (Q_w) from the collection center and transported properly to the production plant. The collected wastes are allowed for inspection first at the production plant. The automated inspection policy is used to maintain the hygiene of laborers. The inspection policy reveals the random rate (v) of unusable wastes due to improper methods of throwing trash outside. Rests $(1 - E[v])$ waste is sent for energy production. Wastes are converted to energy at a conversion rate λ . The produced energies are stored and supplied to the city as on-demand (see <Fig. 1>). The total cost of the production system is minimized with the maximized amount of produced energy (Q_e).

The following assumptions are considered for this study:

- A single-stage production model is considered to produce a single type of renewable energy. A renewable energy production system is contemplated to yield renewable energies from wastes.
- The production system produces energy at a constant rate (p_e) and constant demand (d_e). The production rate is greater than the demand ($P_e > d_e$) that means shortages do not exist in the production system.
- A traditional production system is used with an automated inspection policy (Dey et al., 2021) to separate the unusable wastes. Unusable wastes are detected randomly and disposed of from the system.
- Three types of distributions, uniform, triangular, and beta, are found for random scrap rate.

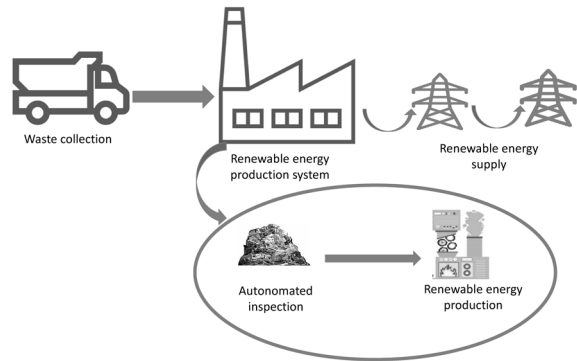


Fig. 1 Renewable Energy Production System from Waste

3. Mathematical Model

In this section, the model formulation and solution procedure are explained elaborately.

3.1 Model Formulation

This study consists of a renewable energy production system. Q_w wastes are transported from the waste center. Those wastes are inspected first with an automated inspection policy, and unusable scraps are discarded from the system with a random rate v . The expected disposed wastes are $E[v]$. Rests $(1 - E[v])$ wastes are sent for energy conversion. The energy conversion rate from waste is λ . Thus, from the Q_w amount of waste $\lambda(1 - E[v])Q_w$ amount of en-

ergy, which is denoted by Q_e is produced. The renewable energy demand d_e of the city is fulfilled from the plant. The costs related to the production plant are explained below.

Raw material collection cost

This study is about renewable energy production systems from waste. Thus, wastes are the raw material collected at a center from different zones of a city. The production plant collects those wastes from the collection center. C_m is the collection cost per ton of the waste for the plant. Hence, the raw material collection cost (RMCC) per cycle is obtained as

$$CC = \frac{C_m Q_w}{T_e} = \frac{C_m d_e}{\lambda(1-E[v])} \quad (1)$$

Transportation and carbon emission cost for raw material collection

Wastes are transported from the collection center by suitable waste transportation trucks. To avoid the spreading of pollution, proper precautions are taken to transport wastes. The distance between the collection center and the production plant is κ . The transportation cost per ton per unit distance is χ_1 . Therefore, the transportation cost (TRP) per cycle can be calculated as

$$TRP = \frac{\chi_1 \kappa Q_w}{T_e} = \frac{\chi_1 \kappa d_e}{\lambda(1-E[v])}$$

The carbon emissions are held during the transportation of wastes. χ_2 is the carbon emissions cost per ton per unit distance. Thus, the carbon emissions cost (CRT) per cycle is

$$CRT = \frac{\chi_2 \kappa Q_w}{T_e} = \frac{\chi_2 \kappa d_e}{\lambda(1-E[v])}$$

Hence, the total transportation and carbon emissions cost (TRCBC) per cycle is

$$TRCBC = \frac{(\chi_1 + \chi_2) \kappa d_e}{\lambda(1-E[v])} \quad (2)$$

Inspection cost of raw material

An automated inspection policy is used to inspect all wastes before conversion. An automated machine does this

type of inspection. As a result, the hygiene of human laborers can be maintained. α_1 is the per ton inspection cost. Therefore, the inspection cost (ICR) per cycle can be written as

$$ICR = \frac{\alpha_1 Q_w}{T_e} = \frac{\alpha_1 d_e}{\lambda(1-E[v])} \quad (3)$$

Raw material scrap cost

After inspection, the random unusable wastes are detected as v . Therefore, $(1-v)$ of the received quantities is the amount of waste used to produce energy. If C_r be the per ton scrap cost, then the expected raw material scrap cost (RMSC) per cycle can be calculated as follows:

$$RMSC = \frac{C_r E[v] Q_w}{T_e} = \frac{C_r E[v] d_e}{\lambda(1-E[v])} \quad (4)$$

Setup cost of the production system

A renewable energy production system is considered to produce renewable energy. The production system needs some initial setup before starting production. The production plant expands the cost S_e per setup as setup cost. Therefore, the setup cost (SCPS) per cycle can be calculated as

$$SCPS = \frac{S_e}{T_e} = \frac{S_e d_e}{Q_e} \quad (5)$$

Holding cost of energy

After producing renewable energy, it is very sensitive to storing energy properly. The average inventory of the production system can be written as $\frac{Q}{2}(1 - \frac{d_e}{P_e})$. C_h is the per unit time storage cost for energy. Thus, the holding cost for energy per cycle is

$$HCE = C_h \frac{Q}{2} (1 - \frac{d_e}{P_e}) \quad (6)$$

Renewable energy supply cost

The renewable energy production system fulfills the demand for energy of a city. The energy should be supplied accurately that customers can get the service properly. Thus, the supply cost is quite crucial. This cost depends on the amount of energy produced. The cost per unit amount of energy is C_s . Therefore, the renewable energy supply cost (ESC) per cycle can be calculated as

$$ESC = \frac{C_s Q_e}{T_e} = C_s d_e \quad (7)$$

Expected total cost per cycle

The total cost can be calculated by summing up all costs of the production system explained above. Thus, the expected total cost per cycle of the renewable production system can be expressed as

$$TC_e = \frac{C_m d_e}{\lambda(1-E[v])} + \frac{(\chi_1 + \chi_2) \kappa d_e}{\lambda(1-E[v])} + \frac{\alpha_1 d_e}{\lambda(1-E[v])} + \frac{C_r E[v] d_e}{\lambda(1-E[v])} + \frac{S_e d_e}{Q_e} + C_h \frac{Q}{2} \left(1 - \frac{d_e}{P_e}\right) + C_s d_e \quad (8)$$

3.2 Solution Methodology

The proposed model is solved analytically by a classical optimization method. Taking the first-order derivative of Equation (8) concerning the decision variable Q_e , it can be obtained as

$$\frac{dTC_e}{dQ_e} = \frac{C_h}{2} \left(1 - \frac{d_e}{P_e}\right) - \frac{S_e d_e}{Q_e^2} \quad (9)$$

The necessary optimization condition can obtain the optimum value of the decision variable. Equating (9) equates to zero to obtain the optimum value of Q_e as

$$Q_e = \sqrt{\frac{2S_e d_e}{C_h \left(1 - \frac{d_e}{P_e}\right)}} \quad (10)$$

The optimum solution can be called the global optimum if the sufficient condition is satisfied. Taking the second-order derivative of Equation (8), it can be written as

$$\frac{d^2 TC_e}{dQ_e^2} = \frac{2S_e d_e}{Q_e^3} > 0 \quad (11)$$

Thus, it can be confirmed that the expected total cost is the global minimum at the optimum Q_e .

The scrap rate is considered a random variable. Based on the model of Sarkar and Sarkar (2020), three different distributions are considered to find the expected random scrap

wastes. The probability distribution functions for different distributions are explained below.

The collected wastes are inspected with the automation policy. The scrap rate is considered as random, which follows a uniform distribution. Thus, the value of $E[v]$ can be obtained as

$$E[v] = \int_u^v x f(x) dx = \int_u^v \frac{x}{v-u} dx = \frac{u+v}{2}$$

where the probability density function (pdf) is

$$f(x) = f(x) = \begin{cases} \frac{1}{v-u}, & u \leq x \leq v \\ 0, & \text{otherwise} \end{cases}$$

If the random scrap rate follows a triangular distribution, the value of $E[v]$ can be written as

$$E[v] = \frac{u+v+w}{3}$$

where the pdf is

$$f(x) = \begin{cases} 0, & x < u \\ \frac{2(x-u)}{(v-u)(w-u)}, & u \leq x < w \\ \frac{2}{v-u}, & x = w \\ \frac{2(v-x)}{(v-u)(v-w)}, & w < x \leq v \\ 0, & v < x \end{cases}$$

If the random scrap rate follows beta distribution, the expected scrapped waste $E[v]$ can be written as $E[v] = \frac{u}{u+v}$, where the pdf of the beta distribution is

$$f(x) = \frac{x^{u-1}(1-x)^{v-1}}{B(u,v)} u > 0, v > 0, 0 \leq x \leq 1.$$

Remark Sarkar and Sarkar (2020) introduced random scrap production rates in their multi-stage biofuel production system. They studied three different types of distributions for random defective items during production, which were scrapped from the system. They compared among the distributions to find the best result. The concept of random scrap rate with different distribution is taken for this study to sort out the wastes for renewable energy production. Both studies are totally different as Sarkar and Sarkar (2020) considered biofuel production, whereas this study regarded renewable energy production from waste. They thought of a multi-stage production system, whereas this study is a single-stage energy production system.

The model's validity can be proved through the numerical experiment in the next section.

4. Numerical Experiment

The numerical experiment has been done by Mathematica 11.3. The data are taken from Sarkar and Sarkar (2020) and Sarkar and Seo (2021). This section has three examples of three types of distributions for scrap rate. The data for those experiments are enlisted in <Table 1>.

Table 1. Data Used in the Numerical Experiments

$d_e = 600\text{MWh/month}$	$C_m = \$2/\text{ton}$
$\chi_1 = \$0.44/\text{ton/Km}$	$\chi_2 = \$0.04/\text{ton/Km}$
$\alpha_1 = \$0.05/\text{ton}$	$C_r = \$0.02/\text{ton}$
$\lambda = 0.54\text{MWh/ton}$	$S_e = \$200/\text{setup}$
$C_h = \$12/\text{MWh/month}$	$P_e = 1000\text{MWh/month}$
$C_s = \$3/\text{MWh}$	$\kappa = 50\text{Km}$

4.1 Numerical Examples

Examples 1, 2, and 3 are explained here based on the different probability distribution of scrap rate of the waste.

Example 1 The random scrap rate follows a uniform distribution

The interval of uniform distribution is considered as $[u, v] = [0.005, 0.250]$. The optimum renewable energy production is $Q_e = 223.61\text{MWh/month}$, and the optimum total cost is $TC_e = \$36050.71/\text{month}$. The graphical representation of the optimum value of TC_e for optimum Q_e is represented by <Fig. 2>.

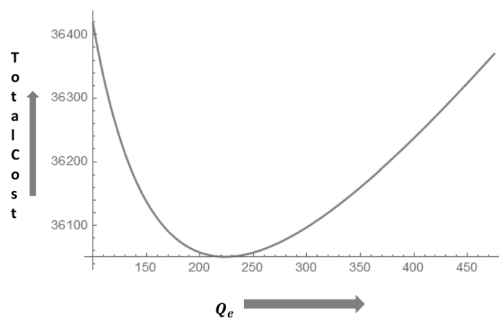


Fig. 2 Optimum Total Cost and Renewable Energy Production for Uniform Distribution

Example 2 The random scrap rate follows a triangular distribution

The lower and upper limit of the triangular distribution is considered the same as the uniform distribution, and the mode of the triangular distribution is $w = 0.015$. Therefore, the optimum value of produced renewable energy is $Q_e = 223.61\text{MWh/month}$, which is the same as the uniform distribution, and the optimum total cost of the production system is $TC_e = \$34682.59/\text{month}$, which is less than Example 1. <Fig. 3> represents the figure of optimum value of TC_e for optimum Q_e .

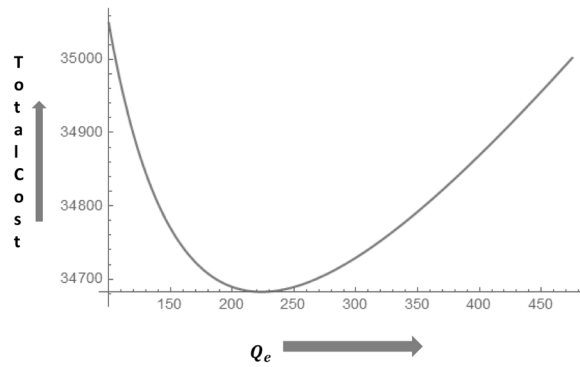


Fig. 3 Optimum Total Cost and Renewable Energy Production for Triangular Distribution

Example 3 The random scrap rate follows a beta distribution

The limit of beta distribution for the random scrap rate is $[u, v] = [0.005, 0.250]$, which is the same as the other two distributions. Thus, the optimum value of produced renewable energy is $Q_e = 223.61\text{MWh/month}$, and the optimum total cost is obtained as $TC_e = \$32397.09/\text{month}$ (see <Fig. 4>), the lowest total cost among the three distributions.

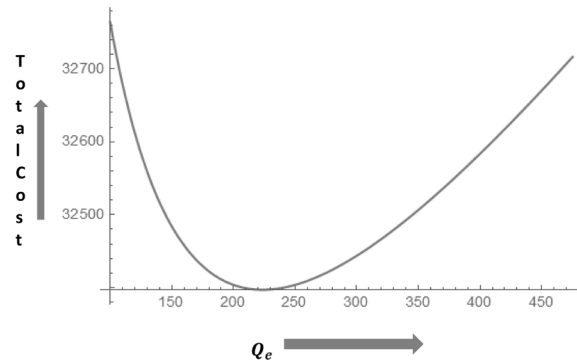


Fig. 4 Optimum Total Cost and Renewable Energy Production for Beta Distribution

4.2 Sensitivity Analysis

The study obtains the lowest cost for beta distribution for the renewable energy production system. Thus, the total cost changes concerning cost parameters have been enlisted in <Table 2>. The changes are investigated for -50%, -25%, 25%, and 50% of the cost parameters. <Fig. 5> shows the changes in expected total cost.

Table 2. Changes of the Expected Total Cost for Changing of Cost Parameters

Parameters	-50%	-25%	25%	50%
C_m	-3.49	-1.75	1.75	3.49
χ_1	-38.44	-19.22	19.22	38.44
χ_2	-3.49	-1.75	1.75	3.49
α_1	-0.09	-0.04	0.04	0.09
C_r	-0.0007	-0.0003	0.0003	0.0007
S_e	-0.98	-0.45	0.39	0.75
C_h	-0.98	-0.45	0.39	0.75
C_s	-2.80	-1.40	1.40	2.80

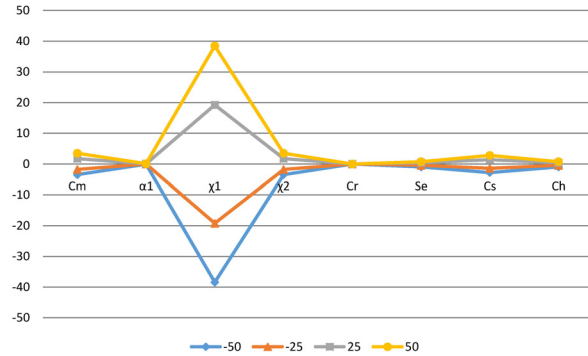


Fig. 5 Changes of the Expected Total Cost for Changing Key Cost Parameters

Based on the sensitiveness of the parameters, the following comments can be drawn for the study.

- The total cost decreases almost 3.50% with a 50% decrease of raw material collection cost, whereas for a 25% decrease of raw material cost, the total cost decreases 1.75%. The same number of changes happens in the reverse direction. Observing the changes in both directions, it can be mentioned that raw material collection cost is a sensitive parameter for the renewable energy production system. The energy plant should take care of

the collection process of waste.

- As inspection of raw material for the production plant has an important role. However, the parameter is not too much sensitive. However, the total cost changes equivalently in both directions with the changes of inspection cost.
- The transportation cost is the most sensitive cost parameter among all cost parameters. The total cost of the energy plant changes 38.44% at a positive and negative direction for 50% changes on both sides. TC_e increases 19.22% for a 25% increase in transportation cost and vice versa. From this result, it can be recommended that waste transportation should be handled very carefully.
- The carbon emission cost for transportation is another sensitive cost parameter. It is exactly sensitive as raw material collection cost. Keep an eye on the sensitiveness of carbon emission. The renewable energy production industry should consider carbon emission reduction for economic and social improvement.
- The third and fourth sensitive parameters are the energy supply and holding costs. The total cost fluctuates almost 3% in both negative and positive directions for changing supply cost, whereas for holding cost, it fluctuates 0.98% in a negative direction and 0.75% in a positive direction. The supply and holding process of energy should be taken care of properly.
- The least sensitive parameter is the scrap cost among all cost parameters. The energy plant can send those unusable wastes for land filling but avoid environmental pollution.

5. Conclusions

The proposed study developed a renewable energy production system, producing energy from a city's wastes. The process could fulfill the energy demand and the city's waste management. A traditional production system was utilized with a fixed production rate. The aim was to reduce the total cost of the energy production system by reducing pollution from wastes.

Even though several resources are for producing renewable energy, wastes were also used to save the environment. The total cost of the plant was minimized at the optimum value

of energy production, and it was found as the global minimum cost. The study was solved analytically and tested numerically with three types of a probability distribution, uniform distribution, triangular distribution, and beta distribution. The automated inspection process found a random rate of scrap wastes.

The global minimum cost was obtained from the theoretical derivation for the closed-form solution of the optimum produced renewable energy. The closed-form of the optimum energy satisfied the basic optimum production formula. Thus, the model was theoretically verified. The numerical study showed that the optimum renewable produced energy for all three distributions was the same, but their corresponding total costs differed. For beta distribution, the total cost was found to be the minimum. For deciding the most important and least parameter among all parameters, it was found that scrap cost is the least effective parameter among all parameters through sensitivity analysis. Even though three separate distributions were considered for the random scrap rate, the corresponding cost was not sensitive. Therefore, the management could manage the cost by the necessity of the cost randomly or constantly, as it did not impact the total cost. The sensitivity analysis revealed the importance of transportation cost, which was the most sensitive cost parameter among all cost parameters.

The study can be extended considering variable transportation costs for waste. The carbon emissions reduction investment can be another extension of this model. An important extension of this study is to consider an automated smart production system controlled by the robotic system. Then the optimum production rate with the best way of inspection through the robotic system. As it was found that the transformation cost was the most sensitive cost, thus an investment can be utilized to control the transportation cost.

REFERENCES

- [1] Ahmed, W. & Sarkar, B. (2019). Management of next-generation energy using a triple bottom line approach under a supply chain framework. *Resources, Conservation and Recycling*, 150, 104431.
- [2] Bhuniya, S., Pareek, S., Sarkar, B., & Sett, B. K. (2021). A Smart Production Process for the Optimum Energy Consumption with Maintenance Policy under a Supply Chain Management. *Process*, 9(1), 19.
- [3] Danish, M. S. S., Senjyu, T., Zaheb, H., Sabory, N. R., Ibrahimi, A. M., & Matayoshi, H. (2019). A novel transdisciplinary paradigm for municipal solid waste to energy. *Journal of Cleaner Production*, 233, 880-892.
- [4] Dey, B. K., Pareek, S., Tayyab, M., & Sarkar, B. (2021). Automation policy to control work-in- process inventory in a smart production system. *International Journal of Production Research*, 59(4), 1258-1280.
- [5] Ducharme, C. (2010). Technical and economic analysis of plasma-assisted waste-to-energy processes research paper I. *School of Engineering and Applied Science*, Columbia University.
- [6] Habib, M. S., Asghar, O., Hussain, A., Imran, M., Mughal, M. P., & Sarkar, B. (2021). A robust possibilistic programming approach toward animal fat-based biodiesel supply chain network design under uncertain environment. *Journal of Cleaner Production*, 278, 122403.
- [7] Lee, J. H. & Moon, I. (2014). Modeling and optimization of energy efficient routing in wireless sensor networks. *Applied Mathematical Modelling*, 38(7-8), 2280-2289.
- [8] Leeuwen, L. B. V., Cappon, H. J., & Keesman, K. J. (2021). Urban bio-waste as a flexible source of electricity in a fully renewable energy system. *Biomass and Bioenergy*, 145, 105931.
- [9] Ramos, A. & Rouba, A. (2020). Renewable energy from solid waste: Life cycle analysis and social welfare. *Environmental Impact Assessment Review*, 85, 106469.
- [10] Rasti-Barzoki, M. & Moon, I. (2020). A game theoretic approach for car pricing and its energy efficiency level versus governmental sustainability goals by considering rebound effect: A case study of South Korea. *Applied Energy*, 271, 115196.
- [11] Rasti-Barzoki, M. & Moon, I. (2021). A game theoretic approach for analyzing electric and gasoline-based vehicles' competition in a supply chain under government sustainable strategies: A case study of South Korea. *Renewable and Sustainable Energy Reviews*, 146, 111139.
- [12] Safarzadeh, S. & Rasti-Barzoki, M. (2019a). A game theoretic approach for assessing residential energy-efficiency program considering rebound, consumer behavior, and government policies. *Applied Energy*, 233-234, 44-61.
- [13] Safarzadeh, S. & Rasti-Barzoki, M. (2019b). A game theoretic approach for pricing policies in a duopolistic

- supply chain considering energy productivity, industrial rebound effect, and government policies. *Energy*, Vol. 167, pp. 92-105.
- [14] Saha, S., Chatterjee, D., & Sarkar, B. (2021). The ramification of dynamic investment on the promotion and preservation technology for inventory management through a modified flower pollination algorithm. *Journal of Retailing and Consumer Services*, 58, 102326.
- [15] Sarkar, B. (2019). Mathematical and analytical approach for the management of defective items in a multi-stage production system. *Journal of Cleaner Production*, 218, 896-919.
- [16] Sarkar, B., Mridha, B., Pareek, S., Sarkar, M., & Thangavelu, L. (2021). A flexible biofuel and bioenergy production system with transportation disruption under a sustainable supply chain network. *Journal of Cleaner Production*, 317, 128079.
- [17] Sarkar, M. & Sarkar, B. (2020). How does an industry reduce waste and consumed energy within a multi-stage smart sustainable biofuel production system? *Journal of Cleaner Production*, 262, 121200.
- [18] Sarkar, M. & Seo, Y. W. (2021). Renewable energy supply chain management with flexibility and automation in a production system. *Journal of Cleaner Production*, 324, 129149.
- [19] Ullah, M., Asghar, I., Zahid, M., Omair, M., Alarjani, A., & Sarkar, B. (2021). Ramification of remanufacturing in a sustainable three-echelon closed-loop supply chain management for returnable products. *Journal of Cleaner Production*, 290, 125609.
- [20] Vandana, Singh, S. R., Yadav, D., Sarkar, B., & Sarkar, M. (2021). Impact of energy and carbon emission of a supply chain management with two-level trade-credit policy. *Energies*, 14(6), 1569.



Mitali Sarkar

Master in mathematics from Jadavpur University, India;
Ph.D from Hanyang University in Industrial Engineering;
Post-doctoral researcher in Seoul Nation University, Yonsei University, and Hanyang University;

Currently working as researcher in Chung-Ang University;
Research area: Supply chain management, Renewable energy, Optimization, Smart manufacturing



서 용 원

서울대학교 산업공학 학사
서울대학교 산업공학 석사
서울대학교 산업공학 박사
현재: 중앙대학교 경영학부 교수
관심 분야: Supply chain management,
Behavioral operations,
Business Analytics

데이터기반 재고관리를 위한 DDPG기반 심층강화학습 알고리즘 연구*

이병권** · 박건수** · 정세윤***†

서울대학교 산업공학과, *한국방송통신대학교 프라임칼리지 첨단공학부

A Modified Deep Deterministic Policy Gradient Algorithm for Data-Driven Inventory Management*

Byeongkwon Lee** · Kun Soo Park** · Se-Youn Jung***†

**Department of Industrial Engineering, Seoul National University, Korea

***Division of Advanced Engineering, Prime College, Korea National Open University, Korea

We develop a modified version of deep deterministic policy gradient(DDPG) algorithm, which is one of the most popular deep reinforcement learning algorithms dealing with continuous action spaces, and apply it to the infinite-horizon newsvendor problems with constant lead time. Reflecting the key features of the inventory management problems, our algorithm, named as Inventory based DDPG (IDDPG), is differentiated in the state variables, action spaces, and cost functions compared to the DDPG. By conducting numerical experiments and comparing the performances, we found that when applied to the inventory management problems, IDDPG is able to find the near-optimal solutions and outperforms the DDPG in most cases. We also found that our IDDPG can be applied to a inventory management problem whose optimal solution is not known.

Keyword : Newsvendor Problem, Data-Driven Approach, Deep Reinforcement Learning, Deep Deterministic Policy Gradient

* This research was supported by Korea National Open University Research Fund

† **Corresponding Author** : Division of Advanced Engineering, Prime College, Korea National Open University, 86 Daehak-ro, Jongno-gu, Seoul 03087, Korea,
Tel: +82-2-3668-4448, E-mail: jseyoun@knou.ac.kr

Received: 21 November 2021, **Accepted**: 10 December 2021

1. 서 론

데이터의 저장 및 연산 성능이 향상되어 데이터의 활용 가치가 높아지면서 재고관리 분야에서도 데이터 기반 의사결정 방법에 대한 연구들이 활발히 진행되어 왔다. 뉴스벤더 문제에서 데이터의 수가 증가함에 따라 최적해에 수렴함을 증명한 Levi et al.(2007)의 Sample Average Approximation(SAA) 알고리즘 연구를 비롯하여 다양한 연구들에서 Empirical Risk Minimization(ERM), Kernel-Weights Optimization(KO) 등을 다루었다(Levi et al., 2007; Levi et al., 2015; Ban & Rudin, 2019). 그러나, 기존의 연구방식은 최적해를 구할 수 있는 재고관리 문제에만 적용가능하기 때문에 확장성이 낮다는 문제점이 존재한다. 이러한 문제의 해결책으로 다양한 산업 환경에서 좋은 성과를 보이고 있는 심층강화학습(Deep Reinforcement Learning, Deep RL)기반 알고리즘이 떠오르고 있다. Gijbrecchts et al.(2020)은 심층강화학습이 유실 판매(lost sales), 복수 조달(dual sourcing), 다체계(multi-echelon) 등을 포함하는 재고관리 문제로도 확장가능함을 보였다. 그러나, 현재까지 진행된 심층강화학습을 적용한 재고관리 연구들은 작은 크기의 이산 행동 공간을 다루거나 연속 행동 공간을 제한적으로 다룬다는 한계가 있다. 또한, 재고관리 문제의 특성을 고려하지 않고 심층강화학습을 단순 적용하는 단계에 머무른 연구들이 종종 발견된다.

전통적인 재고관리 연구들은 문제를 해결하기 위해 수요 분포 추정과 주문량 최적화 결정을 분리시켜 순차적으로 진행하는 Separated Estimation and Optimization(SEO) 접근법을 사용한다. 또한, 복잡한 현실문제를 간단하게 모델링하기 위해 수요 분포에 대해서는 독립항등분포(independent and identically distributed random variables) 가정과 지수족(exponential family) 가정들이 흔히 사용된다. 그러나, 연구에 사용된 가정들로 이루어진 상황이 현실과 다른 경우에는 연구결과로 제시된 성능을 보장하지 못하는 경우가 많다. 특히, SEO 접근법을 사용한 재고관리 의사결정은 수요 분포 추정 결과에 강하게 의존하여 주문량 최적화 결정이 이루어지기 때문에 수요 분포에 대한 가정이 현실과 괴리가 클수록 성능이 저하된다. 이러한 문제점을 해결하기 위해 수요 분포를 정확히 알고 있다고 가정하는 것이 아니라 현실의 상황과 비슷하게 과거의 수요 데이터를 가지고 있다는 가정을 사용하기 시작했다(Ban & Rudin, 2019). 또한, 데이터를 기반으로 추정한 수요 분포가 정확하다고 믿고

주문량 최적화 결정을 진행하는 것이 아니라 추정한 수요 분포에 대한 불확실성 정도를 고려하는 방향으로의 연구가 진행되고 있다(Prak & Teunter, 2019).

SEO 접근법은 광범위하게 사용됨에도 불구하고 두 단계에 걸쳐서 의사결정을 진행하는 방식의 최적해가 문제 전체의 최적해임을 보장하지 못한다는 단점이 있다(Liyanage & Shanthikumar, 2005). 이런 문제점을 극복하기 위해 Joint Estimation Optimization(JEO) 접근법에 대한 연구가 최근에 관심을 받고있다. JEO 접근법은 수요를 예측하는 과정을 중간에 거치지 않고 특성 데이터(feature data)에 의존하여 주문량을 결정하는 단일 단계의 의사결정 방법을 의미한다. 대표적으로 비모수적 데이터기반 휴리스틱 접근법인 SAA 알고리즘이 있다(Homem-De-Mello, 2000; Kleywegt et al., 2001; Levi et al., 2007; Levi et al., 2015). 다수의 특성 데이터를 활용한 알고리즘으로는 ERM 알고리즘과 KO 알고리즘이 있다(Ban & Rudin, 2019).

현재까지 대부분의 JEO 접근법 연구들은 수요 분포를 알고 있는 상태에서 SEO 접근법으로 최적해를 구할 수 있는 경우에 대해서만 다뤘었다. 앞서 언급한 SAA, ERM, KO 모두 뉴스벤더 문제의 최적해가 재고부족비용(underage cost)과 재고초과비용(overage cost)의 긴급률(critical ratio)에 의해 결정된다는 특성 때문에 수요 분포 추정에 해당하는 최적화 문제만을 풀면 되는 구조를 가지고 있다. 그러나, 이러한 JEO 접근법 연구들은 SEO 접근법으로 최적해를 도출하지 못하는 문제들에 대해서는 확장가능성이 낮다는 문제가 있다.

심층강화학습 알고리즘은 한 가지 문제에 국한되지 않고 다양한 문제에 적용가능하며, 심지어 성능도 기존에 존재하던 방식들에 비해 월등히 뛰어난 모습을 보여 주기도 한다. 실제로 로봇틱스, 자율 주행, 추천, 얼굴 인식, 데이터 센터 쿨링, 주식 거래, 질병 진단, 자연어 처리, 프로그램 작성, 수학 증명 등의 광범위한 시스템을 위해 연구 및 활용되어 심층강화학습의 효용성을 입증하고 있다(Li, 2018). 이처럼 많은 분야에서 좋은 모습을 보여주고 있기 때문에 재고관리 분야에서도 심층강화학습 적용 연구에 대한 관심이 생기고 있다. 특히, JEO 접근법 이면서 다양한 재고관리 문제로의 확장성이 높다는 장점 때문에 재고관리 분야의 심층강화학습 연구가 활발히 이루어지고 있다(Gijbrecchts et al., 2020).

현재까지 최적해를 구할 수 있는 알고리즘이 개발되지 않은 재고관리 문제들에 심층강화학습을 적용시킨 연구들이 이루어졌다. 예를 들어, 비어 게임(beer game),

유실 판매, 복수 조달, 다체계 등의 문제들에서 심층강화학습으로 근사 최적 정책(near-optimal policy)을 찾는 모습을 볼 수 있다. 특히, Oroojlooyjadid et al.(2020) 연구는 전통적인 방식으로 연구하기 까다로운 조건의 비어 게임 상황에서 좋은 성과를 보이는 의사결정 정책을 구할 수 있음을 보여준다. 그러나, 심층강화학습이 좋은 성능을 보인다는 하더라도 수렴성에 대한 이론적인 증명이 이루어지지 않은 상태이기 때문에 단순한 수치실험의 성능으로 알고리즘을 신뢰하기 힘들 수 있다(Fan et al., 2019). 또한, 심층강화학습을 적용한 재고관리 연구들이 모두 최적해와의 비교가 힘든 상황에서 휴리스틱과 상대적인 성능 평가로 이루어졌기 때문에 분석할 수 있는 부분이 제한적이라는 문제가 있다. 따라서, 본 연구에서는 최적해가 알려진 뉴스벤더 문제에서 심층강화학습 알고리즘의 해와 최적해를 비교하여 알고리즘의 성능을 객관적이고 체계적으로 검증하고자 한다.

본 연구는 연속 행동 공간을 다루는 심층강화학습 알고리즘 중 대표적인 Deep Deterministic Policy Gradient(DDPG)를 재고관리 문제에 적용하여 성능을 검증하고자 한다. 특히, 본 연구에서는 재고관리 문제에 DDPG를 단순 적용하는 것이 아니라, 재고관리 문제의 특성들을 고려하여 상태 변수(state variable), 행동 공간(state space), 비용 함수 등을 적절히 변형시킨 Inventory-based DDPG(IDDPG) 알고리즘을 제안한다. 심층강화학습 분야의 특성상 수리적인 분석을 통해 최적 여부를 판단하는 것이 어렵기 때문에 최적해를 구할 수 있는 재고관리 문제 중 리드타임이 일정한 무한 기간 뉴스벤더 문제(infinite horizon newsvendor problem with constant lead time)를 다룬다. 수치실험을 진행하여 DDPG와 IDDPG가 도출하는 해들을 최적해와 비교하여 성능을 측정하고 분석하고자 한다. 이를 통해, IDDPG가 DDPG 대비 충분히 좋은 해를 도출할 수 있을 뿐만 아니라, 최적해에도 근사한 해를 도출할 수 있음을 입증하고자 한다. 마지막으로, 본 연구에서는 최적해가 알려져 있지 않은 비정형적이고 복잡한 재고 문제로도 IDDPG를 확장 가능함을 간단한 실험 예제를 통해 검증하였다.

2. 선행 연구

본 장에서는 연구와 관련된 선행 연구를 데이터기반 재고관리, 심층강화학습, 그리고 심층강화학습을 활용한 재고관리 세 가지로 분류하여 조사하고 비교하였다.

2.1 데이터기반 재고관리 연구

데이터기반 재고관리 연구는 특정한 수요의 분포를 가정하는 대신 관측한 수요 데이터를 기반으로 하는 재고 의사결정을 다룬다. Levi et al.(2007)는 수요가 독립항등분포를 따른다는 가정 하에서의 샘플 데이터만 사용 가능한 상황에서 뉴스벤더 문제를 푸는 알고리즘으로 SAA를 제안하였다. 이 때, 샘플 데이터를 N 개 가지고 있는 상태에서 SAA 알고리즘으로 도출한 해가 epsilon optimal일 확률의 하한을 제시했다. Levi et al.(2015)에서는 Levi et al.(2007) 연구에서 제시했던 하한보다 좋은 하한을 제시했다. Ban & Rudin(2019)에서는 수요가 단일 분포에서 발생하는 것이 아니라 특성 데이터에 따라 다른 분포에서 발생하는 것이라면 SAA 알고리즘이 최적해에 수렴하지 못함을 보였다. 이런 문제점에 대한 해결책으로 ERM 알고리즘과 KO 알고리즘을 제시했다. 그러나, Prak & Teunter(2019)에서는 SAA, ERM, KO와 같은 알고리즘은 샘플 데이터의 수가 증가할수록 최적해에 수렴하더라도, 표본집단과 모집단 사이의 차이에서 비롯된 불확실성을 고려하지 않으면 샘플 데이터의 수가 적은 경우 비용의 기댓값에 대한 성능이 나쁘다는 것을 보여준다. 수렴성과 epsilon optimal일 확률의 하한과 같이 분석적인 측면에서 훌륭한 알고리즘을 제시하는 것도 중요하지만, 실질적으로 중요한 비용의 기댓값 측면에서 훌륭한 알고리즘에 대한 연구가 필요하며, 심층강화학습 알고리즘이 대안이 될 수 있다.

2.2 심층강화학습 연구

심층강화학습은 인공지능망과 같은 딥러닝 기술을 강화학습에 적용시켜 순차적 의사결정 문제를 해결하는 방법론으로서, 마르코프 결정 과정(Markov Decision Process, MDP)으로 문제를 모형화하고 가치 함수(value function) 또는 정책 함수(policy function)에 대한 근사자(approximator)로 인공지능망을 활용하는 것을 의미한다. 심층강화학습은 환경(environment)에 대한 모델의 존재 여부에 따라 model-free 강화학습과 model-based 강화학습으로 구분된다. model-based 강화학습은 행동(action)에 따라서 환경이 어떻게 변화할지 예측이 가능하고 미리 변화를 예상하고 행동을 계획하여 실행할 수 있다는 장점이 있기 때문에 대체로 학습의 효율이 좋다(Kumar et al., 2019). 그러나, 대부분의 문제들은 환경의 정확한 모델을 알아내기 어렵기 때문에 활용도가 광범위

한 model-free 강화학습에 대한 연구가 더 활발히 진행되었다.

심층강화학습은 근사하는 대상이 가치 함수인지 정책 함수인지 혹은 둘 다인지에 따라서 각각 가치 기반 강화학습(value-based RL), 정책 기반 강화학습(policy-based RL), 액터-크리틱 강화학습(actor-critic RL)으로 분류된다. 가치 기반 강화학습에는 딥마인드에서 발표한 Deep Q-Network(DQN)와 Asynchronous Advantage Actor-Critic(A3C) 알고리즘이 있다(Mnih et al., 2015; Mnih et al., 2016). 정책 기반 강화학습에는 기본이 되는 REINFORCE 알고리즘, 우수한 성능을 보여준 trust region policy optimization(TRPO) 알고리즘, TRPO 만큼의 성능을 보이면서 구현이 매우 쉽고 데이터 효율성을 높인 proximal policy optimization(PPO) 알고리즘 등이 있다(Williams, 1992; Schulman et al., 2015; Schulman et al., 2017). 정책 기반 강화학습은 원하는 방향으로 직접적으로 최적화하기 때문에 학습 결과가 안정적이고 신뢰도가 높다. 반면, 가치 기반 강화학습은 self-consistency 방정식을 만족하는 Q 함수를 학습해서 에이전트(agent)의 성능을 간접적으로 최적화하고, 안정성이 떨어지지만 상대적으로 샘플 데이터를 재활용할 수 있는 가능성이 높기 때문에 데이터 효율성이 높다는 장점이 있다. 따라서, 이러한 장단점에 대한 균형을 맞추기 위해 가치 함수와 정책 함수에 대한 근사를 모두 활용하는 방향으로 정책들이 발전하여 액터-크리틱 강화학습인 DDPG와 soft actor-critic(SAC)이 개발되었다(Lillicrap et al., 2016; Haarnoja et al., 2018). 본 연구에서는 결정론적 정책(deterministic policy)과 Q 함수를 서로 발전시켜 동시에 학습시키는 DDPG를 뉴스벤더 문제에 활용하는 방법에 대해 다룬다. 또한, DDPG와는 달리 액터 네트워크와 크리틱 네트워크의 상태 변수를 다르게 설정하여 뉴스벤더 문제의 특성을 활용할 수 있도록 변형시킨 IDDPG 알고리즘을 소개하고자 한다.

2.3 심층강화학습을 적용한 재고관리 연구

최근에 재고관리의 다양한 문제들에 대해 심층강화학습을 적용하는 연구가 활발히 진행되고 있다. 수요가 불확실한 상황에서의 용량 제한이 있는 다기간 공급 사슬 최적화 문제(Peng et al., 2019), 비정상성 수요 하에서의 다제품 재고관리 문제(Suetsugu et al., 2019), 유실 판매 · 복수 조달 · 다체계 재고관리 문제(Gijsbrechts et al., 2020), 비어 게임(Oroojlooyjadid et al., 2020; Dittrich

& Fohlmeister, 2021), 다품목 일괄구매(joint replenishment) 문제(Vanvuchelen et al., 2020), 부패 가능한 제품 재고관리 문제(De Moor et al., 2021) 등 대부분의 연구들이 근사 최적 정책을 구할 수 있음을 보였다. 특히, Oroojlooyjadid et al.(2020)는 기존에 연구되지 않았던 비어 게임 구성원들이 비합리적인 주문을 하는 상황에서 심층강화학습을 통해 좋은 의사결정을 할 수 있음을 보여줬다. 하지만, 대부분의 연구가 작은 이산 행동 공간과 작은 이산 수요 분포를 기반으로 한 수치실험 환경에서 진행되었다. 이는 최적해를 정확하게 구할 수 없는 문제이기 때문에 발생한 문제로 연구결과의 확장성과 신뢰성의 한계가 있다.

재고관리 연구에 사용된 심층강화학습 알고리즘의 종류는 상대적으로 적다. 가치 기반 강화학습으로는 DQN과 branching DQN(BDQN)이 활용되었고, 정책 기반 강화학습으로는 vanilla policy gradient(VPG), PPO가 활용되었다. 액터-크리틱 강화학습은 A3C가 유일하게 활용된 것으로 보인다. DQN, BDQN, A3C, PPO는 이산적인 행동 공간을 다루었으며 VPG만이 연속적인 행동 공간을 다루었다. Peng et al.(2019)에서는 기초적인 정책 기반 알고리즘인 VPG를 사용하여 연속적인 행동 공간을 다루기는 했으나 근사 최적 정책을 찾을 수 있음을 보이지는 못했기 때문에 연속적인 행동 공간을 잘 다루었다고 보기 힘들다. 따라서, 본 연구는 다양한 분야에서 훌륭한 성능을 보였던 액터-크리틱 강화학습인 DDPG를 활용해 연속적인 행동 공간을 다루었다는 점에서 의의가 있다.

아직까지 재고관리 문제가 지니는 특성들을 적절히 고려하여 심층강화학습을 적용한 연구는 충분히 이루어지지 않았으며, 대부분의 연구가 심층강화학습을 재고관리 문제에 단순 적용하는 수준으로 진행되었다(Peng et al., 2019; Gijsbrechts et al., 2020; Vanvuchelen et al., 2020; Dittrich & Fohlmeister, 2021). De Moor et al.(2021)는 DQN의 보상 함수를 변형시켜 부패 가능한 제품을 다루는 재고관리 문제에서 해의 성능과 안정성을 높일 수 있음을 보였지만, 재고관리 문제의 특성을 활용했다고 보기는 힘들다. 재고관리 문제의 특성을 고려하여 성능을 개선시키려는 연구 중 하나로 Suetsugu et al.(2019)는 다제품 재고관리 문제의 특성을 고려하여 BDQN을 기반으로 제품별로 보상을 할당하는 구조를 추가한 BDQN with reward allocation(BDQN-RA)을 제시하였다. Oroojlooyjadid et al. (2020)는 여러 참여자들이 정보 공유가 어려운 상태에서 의사결정을 하는 비어 게임의 특성을 고려하여 적절히 전체 시스템의 이익을 고려할 수 있도록 보상 함수를 변형한 shaped-reward DQN을

제시하였다. 이처럼 기존의 연구들은 보상 함수를 변형하는 것에 초점을 맞춘 것을 볼 수 있다. 본 연구에서는 보상 함수를 변형하는 것뿐만 아니라, 다양한 재고관리 문제들이 보이는 백오더 특성을 활용하여 상태 변수의 차원을 낮추고 행동 공간을 변형시켜 심층 신경망의 학습을 효율적으로 할 수 있는 방법을 제안하고자 한다.

3. 모델

심층강화학습 알고리즘은 동적 계획법(dynamic programming)을 통해 풀기 힘든 큰 규모의 마르코프 결정 과정 모델을 근사하여 풀 수 있도록 도와준다. 따라서 심층강화학습을 적용하기 위해 리드타임이 일정한 무한 기간 뉴스벤더 문제를 마르코프 결정 과정으로 모델링하고자 한다. 본 문제에서는 전통적인 뉴스벤더 문제와 같이 제품 한 개를 주문하는데 단위 비용 c 가 발생하며 제품 도착까지 리드타임 l 만큼의 시간이 발생한다. 또한, 주기(period) $t \geq 1$ 마다 수요 d_t 가 발생하며 남은 재고량만큼 재고유지비용(holding cost) h 가 발생하고 충족시키지 못한 수요량만큼 재고부족비용(shortage cost) p 가 발생한다. 이 때, 수요 d_t 의 분포는 모른다고 가정한다.

주기 t 가 시작할 때 주기 $t-1$ 의 기말재고 I_{t-1} 와 운송 중 재고(outstanding receipts, pipeline vector) 벡터 $Q_{t-1} = (q_{t-l}, q_{t-l+1}, \dots, q_{t-1})$ 를 기반으로 주문량 $q_t \geq 0$ 을 결정한다. 이 때, 운송 중 재고 벡터는 리드타임이 양수인 경우만 고려한다. 기말재고 I_{t-1} 에 주문량 q_{t-l} 를 더한 재고를 가지고 수요 d_t 를 만족시키며, 초과된 수요는 백오더(backorder)되어 향후에 재고가 생겨서 충족될 때까지 유지된다. 백오더된 수요량은 기말재고 I_t 의 값을 음수로 나타내어 표현한다. 따라서, 주기 t 의 기말재고 $I_t = I_{t-1} + q_{t-l} - d_t$ 의 값을 가지고 운송 중 재고 벡터 $Q_t = (q_{t-l+1}, \dots, q_t)$ 가 된다. 뉴스벤더 문제는 주기 t 마다 상태 벡터 $S_t = (I_{t-1}, Q_{t-1})$ 에서 행동 $a_t = q_t$ 을 결정하여 비용을 최소화하는 마르코프 결정 과정으로 모델링될 수 있다.

리드타임 l 이 양수인 경우 재고관리 문제의 특성을 활용하여 상태 벡터 S_t 의 차원을 줄일 수 있다. 주기 t 에서 주문량 q_{t-l} 만큼 공급을 받아 기말재고 $I_{t-1} + q_{t-l}$ 만큼의 재고를 활용할 수 있게 된다. 따라서, 주기 t 에서 기말재고 I_{t-1} 과 주문량 q_{t-l} 의 값을 모두 알고 있는 조건에서 주문량 q_t 를 결정하는 것과 두 값의 합만 알고 있는 조건에서 주문량 q_t 를 결정하는 것 사이에 차이가 없다. 뿐만

아니라, 이러한 정보의 차이가 미래에 발생할 비용의 기댓값을 예측하는데 영향을 주지 않기 때문에 동일한 마르코프 결정 과정이라고 볼 수 있다. 즉, 상태 벡터 S_t 를 ($s_0 = I_{t-1} + q_{t-l}$, $s_1 = q_{t-l+1}$, $s_2 = q_{t-l+2}$, ..., $s_{l-1} = q_{t-1}$)와 같이 정의하여 벡터의 차원을 $l+1$ 에서 $\max\{1, l\}$ 로 줄일 수 있다.

뉴스벤더 문제는 전체 비용 최소화를 목적으로 가능한 모든 상태 S_t 에 대해 어떤 행동 a_t 를 선택할지를 결정하는 정책 함수 π 를 최적화하는 문제로 볼 수 있다. 주기 t 마다 발생하는 비용은 수식 (1)과 같이 주기 t 의 상태 S_t 와 행동 a_t 의 함수로 이루어진다. 무한한 주기에서 발생하는 비용들을 주기 t 의 현재 가치로 계산하여 전체 비용 지표로 상태 가치 함수(state-value function) $v^\pi(S_t)$ 와 상태-행동 가치 함수(action-value function) $Q^\pi(S_t, a_t)$ 를 사용한다. 상태 가치 함수 $v^\pi(S_t)$ 는 상태 S_t 에서 정책 π 를 취해서 주기 t 부터 발생하는 비용들의 현재 가치를 의미하며 수식 (2)와 같다. 상태-행동 가치 함수 $Q^\pi(S_t, a_t)$ 는 상태 S_t 에서 행동 a_t 를 취한 후에 정책 π 를 취해서 발생하는 비용들의 현재 가치를 의미하며 수식 (3)과 같다. 이 때, 상수 $\gamma \in (0, 1)$ 은 할인 계수(discount factor)를 의미하며 E^π 은 정책 π 를 취할 때의 기댓값 연산자를 의미한다. 비용 $c(S_{t+j}, a_{t+j})$ 는 불확실성을 가진 수요에 의해 결정되는 확률변수이기 때문에 기댓값을 지표로 사용한다.

$$c(S_t, a_t) = q_t c + h[I_t]^+ + p[I_t]^- \quad (1)$$

$$v^\pi(S_t) = E^\pi \left[\sum_{j=0}^{\infty} \gamma^j c(S_{t+j}, a_{t+j}) \mid S_t \right] \quad (2)$$

$$Q^\pi(S_t, a_t) = E^\pi \left[\sum_{j=0}^{\infty} \gamma^j c(S_{t+j}, a_{t+j}) \mid S_t, a_t \right] \quad (3)$$

$$v^*(S_t) = \min_{a_t \in A} \{ c(S_t, a_t) \quad (4)$$

$$+ \gamma \sum_{S' \in \mathcal{S}} P(S_{t+1} = S' \mid S_t, a_t) v^*(S') \} \\ Q^*(S_t, a_t) = c(S_t, a_t) \quad (5)$$

$$+ \gamma \sum_{S' \in \mathcal{S}} P(S_{t+1} = S' \mid S_t, a_t) \min_{a' \in A} Q^*(S_{t+1}, a')$$

최적해를 도출하기 위해서는 가능한 정책들의 집합을 정책 집합 Π 에서 모든 상태 S_t 에 대해 $v^\pi(S_t)$ 의 하한 $v^*(S_t)$ 를 구하거나 모든 상태 S_t 와 행동 a_t 에 대해 $Q^\pi(S_t, a_t)$ 의 하한 $Q^*(S_t, a_t)$ 를 구해야 한다. $v^*(S_t)$ 를 구하기 위해서는 수식 (4)의 벨만 방정식을 사용해 가치 반복(value iteration) 알고리즘을 적용하고 $Q^*(S_t, a_t)$ 를 구하기 위해서는 수식 (5)의 벨만 방정식을 사용해 정책

반복(policy iteration) 알고리즘을 적용한다(Bellman, 1954). 이 때, 상태 공간 $\mathcal{S}$ 는 허용 가능한 상태 벡터들의 집합을 의미하고 행동 공간 $\mathcal{A}$ 는 선택 가능한 행동의 집합을 의미한다.

$v^*(S_t)$ 또는 $Q^*(S_t, a_t)$ 가 존재하는 MDP 문제라면 벨만 방정식을 무수히 반복하여 해를 도출할 수 있지만 뉴스벤더 문제처럼 상태 공간, 행동 공간, 수요의 불확실성으로 인한 상태 전이 행렬(state transition matrix)들의 차원에 의해 복잡도가 높은 문제의 경우는 이 방법을 사용하여 해를 도출하는 것이 현실적으로 어렵다. 뿐만 아니라, 수요 d_t 의 분포를 모르는 상황처럼 전이 확률 $P(S_{t+1} = S' | S_t, a_t)$ 를 모르는 경우 벨만 방정식으로 최적해를 구하기 어렵다. 따라서, 최적의 상태-행동 함수 $Q^*(S_t, a_t)$ 또는 최적의 정책 함수 π^* 를 근사하는 강화학습 알고리즘을 활용해야 한다. 이 때, 근사하는 대상에 따라서 분류되는 가치 기반 강화학습, 정책 기반 강화학습, 액터-크리틱 강화학습들이 활용될 수 있다. 본 연구에서는 대표적인 액터-크리틱 강화학습 알고리즘인 DDPG를 통해 뉴스벤더 문제를 다룬다.

4. 알고리즘

본 장에서는 기존의 DDPG를 재고관리 문제에 적용하기 적절하도록 상태, 행동, 비용 등의 구조를 적절하게 변형한 Inventory-based DDPG(IDDPG, Algorithm 1) 알고리즘을 설명한다.

4.1 상태 변수

DDPG에서는 정책 함수를 근사하는 액터 네트워크 $\mu_\theta(S)$ 와 상태-행동 가치 함수를 근사하는 크리틱 네트워크 $Q_\phi(S, a)$ 의 파라미터 θ 와 ϕ 를 학습하여 뉴스벤더 문제를 해결한다. 이 때, 파라미터 θ 와 ϕ 는 심층 신경망(deep neural network)으로 설계된 액터와 크리틱 네트워크의 가중치(weight) 및 편향(bias) 파라미터들의 전체 집합을 의미한다. 네트워크의 파라미터 집합의 크기는 상태 S 의 차원에 비례하게 커지기 때문에 뉴스벤더 문제의 경우 리드타임이 증가할수록 학습의 난이도가 어려워짐을 알 수 있다. 따라서, 본 연구에서는 액터-크리틱 구조의 심층강화학습에서 상태 S 의 차원을 최소화하여 심층 신경망의 크기를 줄일 수 있는 방법을 제안하고자 한다.

리드타임이 일정한 백오더 재고관리 문제 하에서는 기말재고 I_{t-1} 과 운송 중 재고 벡터 Q_{t-1} 의 정보를 사용하는 대신 단일 차원의 숫자인 재고상태(inventory position) $IP_t = I_{t-1} + \sum_{i=t-l}^{t-1} q_i$ 정보만을 활용해도 최적해를 제시할 수 있음을 알고 있다(Arrow et al., 1958). 그러나, 동일한 IP_t 값을 가지면서 기말재고 I_{t-1} 과 운송중 재고 벡터 Q_{t-1} 값이 다른 두 상태 S' 와 S'' 의 상태-행동 가치 함수 $Q^*(S', a)$ 와 $Q^*(S'', a)$ 의 값은 다를 수 있다. 즉, 액터 네트워크에서는 IP_t 를 입력값(input value)으로 사용할 수도 있는 반면 크리틱 네트워크에서는 상태 S 를 입력값으로 사용해야만 한다. 따라서, DDPG가 액터와 크리틱 네트워크에서 동일한 상태 S 를 사용하는 것과 다르게 IDDPG는 액터와 크리틱 네트워크에 상태 $S^A = \sum_{i=0}^{l-1} s_i$ 와 상태 $S^C = (s_0^C = I_{t-1} + q_{t-l}, s_1^C = q_{t-l+1}, \dots, s_{l-1}^C = q_{t-1})$ 를 각각 사용한다. 즉, IDDPG는 액터 네트워크의 상태 변수 차원을 $\max\{1, l\}$ 에서 1로 낮춰 학습의 효율을 높일 수 있다.

4.2 행동 공간

재고관리 문제의 특성상 주문량은 음수가 될 수 없기 때문에 행동 공간의 비대칭성이 존재한다. 즉, 재고상태 IP 가 적정수준보다 낮아질수록 주문량을 늘림으로써 적절한 대응을 할 수 있지만 필요 이상으로 높은 경우에 대해서는 주문을 하지 않는 것 외에는 할 수 있는 것이 없다. 따라서, 주문량 q 가 특정 양수값을 가지는 경우보다 0인 경우의 중요성이 높다. 본 연구에서는 주문량 $q > 0$ 에 대응되는 행동 a 의 범위보다 주문량 $q = 0$ 에 대응되는 행동 a 의 범위를 크게 설정하여 이러한 중요도 차이를 고려하는 방법을 제시하고자 한다.

기존의 심층강화학습을 적용한 재고관리 연구들에서는 크게 두 가지 방식으로 행동 a 를 정의한다. 주문하기 직전의 수요 d 를 기준으로 얼마나 많은 혹은 적은 양을 주문할지 정하는 $d+a$ 규칙을 사용할 수 있다. 이 방식은 유한한 작은 행동 공간 $A = [a_l, a_u]$ 으로 무한한 공간의 주문량 $q \in [0, \infty)$ 을 다룰 수 있다는 장점이 있지만, 직전의 수요 d 에 대한 정보는 상태 S 에 담겨있기 때문에 심층 신경망에 중복된 영향을 줄 수 있다. 또한, 재고상태 IP 가 높아 주문량 $q = 0$ 이 최적인 상황에서 $d+a_l > 0$ 이기 때문에 최적의 행동을 할 수 없는 경우가 발생할 수 있다.

대안으로 행동 a 와 주문량 q 를 동일하게 정의할 수 있다. 이는 유한한 행동 공간 $A = [0, a_u]$ 에 일대일대응으

로 주문량 $q \in [0, a_u]$ 가 결정됨을 의미한다. 이 방식은 (a_u, ∞) 에 해당하는 주문량 q 을 다루지 못한다는 단점이 존재하지만, 주문량의 한계가 존재하는 현실 상황을 반영한다는 점에서 큰 제약은 아니다. 그러나, 여전히 주문량 $q=0$ 에 해당하는 경우가 행동 $a=0$ 하나뿐이라는 단점이 존재한다. 따라서, 본 연구는 행동 공간 $A=[a_l, a_u]$ 에서 행동 $a \in [a_l, 0]$ 을 주문량 $q=0$ 에 해당되도록 행동 공간을 정의하는 방식을 제안한다. 이 때, 수치 실험에서는 주문량의 한계가 존재하지 않는 $a_u = \infty$ 의 경우를 다루기 때문에 a_u 를 충분히 크게 설정한다.

4.3 상태 공간과 비용 함수

상태 변수와 행동 공간을 정의했기 때문에 상태 공간 $\mathcal{S}$ 를 다음과 같이 정의할 수 있다. 상태 s_0^C 는 허용 가능한 최대 백오더량 $\bar{B}$ 과 기말재고량 $\bar{I}$ 에 따라서 $s_0^C \in [-\bar{B}, \bar{I}]$ 을 만족하고, 상태 $s_1^C, \dots, s_{l-1}^C$ 는 과거의 주문량 q 에 대한 기록이기 때문에 $s_1^C, \dots, s_{l-1}^C \in [0, a_u]$ 를 만족하며, 상태 S^A 는 재고상태의 상한과 하한을 의미하기 때문에 $S^A \in [-\bar{B}, \bar{I} + l^* a_u]$ 을 만족한다. 본 연구에서는 백오더와 재고의 용량제한이 없는 경우에 대해서 다루기 때문에 $\bar{B}$ 와 $\bar{I}$ 를 충분히 크게 설정하여 실험한다.

현실에 적용하는 과정에서는 상태 공간 $\mathcal{S}$ 를 벗어나는 경우가 없음에도 불구하고, 심층강화학습을 학습하는 과정에서는 발생할 수 있다. 보통 학습과정에서 상태 공간 $\mathcal{S}$ 를 벗어나면 기존과 동일한 비용함수만큼의 비용이 발생하고 에피소드를 종료하는 조치를 취한다. 상태 공간 $\mathcal{S}$ 를 크게 설정했기 때문에 재고관리 문제의 특성상 최적의 의사결정과 비슷하게 행동하면 상태 공간 $\mathcal{S}$ 의 경계에 가까운 상태가 발생하기 굉장히 힘들다. 이러한 특성 때문에 상태 공간 $\mathcal{S}$ 의 경계에서의 의사결정을 학습할 기회가 적은 문제가 발생한다. 본 연구에서는 이러한 문제를 완화시키기 위해 상태 공간 $\mathcal{S}$ 를 벗어나는 경우 패널티 K 를 발생시켜 비용함수가 $c(S, a) + K$ 이 되도록 하는 방식을 적용하였다.

4.4 IDDPG 알고리즘

IDDPG는 결정론적으로 상태 S^A 와 행동 a 를 매핑하는 정책 함수를 근사하는 액터 네트워크 $\mu_\theta(S^A)$ 와 상태 S^C 와 행동 a 의 상태-행동 가치 함수를 근사하는 크리틱 네트워크 $Q_\phi(S^C, a)$ 를 구성하는 매개변수 θ 와 ϕ 를 재고

관리 문제에 적합하게 학습시키는 것을 목표로 한다.

연속 행동 공간을 다루는 문제에서는 Q-러닝의 벨만 방정식을 그대로 활용하기 힘들다. 매 스텝마다 모든 행동 a 에 대해 최적화하기 위해서는 많은 시간이 소요되기 때문이다. 따라서, 크리틱 네트워크는 수식 (6)을 목표 값(target value)으로 설정하여 벨만 방정식을 근사한다. 이 목표 값을 사용하여 수식 (7)과 같은 손실함수를 최소화하는 것을 목표로 심층 신경망을 학습한다. 액터 네트워크는 Silver et al.(2014)에서 증명된 정책 경사(policy gradient) 방식인 수식 (8)을 통해 학습한다.

수식 (6)의 목표 값을 계산할 때는 학습하는 네트워크와 다른 목표 네트워크를 사용한다. 이 때, DDPG에서와 동일하게 Mnih et al.(2013) 연구에서 사용된 목표 네트워크를 변형한 완화적 목표 네트워크를 활용한다. 기본적인 목표 네트워크 개념대로 액터와 크리틱 네트워크를 복제하여 목표 네트워크 $\mu_{\theta^-}(S^A)$, $Q_{\phi^-}(S^C, a)$ 를 생성하고 이를 통해 목표 값(target value)을 계산한다. 목표 네트워크들의 가중치(weight) 및 편향(bias) 파라미터들을 업데이트할 때, 학습된 네트워크의 파라미터 θ 와 ϕ 으로 직접적으로 업데이트하는 것이 아니라 완화적으로 진행한다. $\theta^- \leftarrow \tau\theta + (1-\tau)\theta^-$, $\phi^- \leftarrow \tau\phi + (1-\tau)\phi^-$, $\tau \ll 1$. 이는 학습과정에서 목표 값이 빠르게 변화하는 것을 억제시키기 때문에 학습의 안정성을 향상시키는 역할을 한다.

$$y_i^- = c_i + \gamma Q_{\phi^-}(S_{i+1}^C, \mu_{\theta^-}(S_{i+1}^A)) \quad (6)$$

$$L = \frac{1}{N} \sum_i (Q_\phi(S_i^C, a_i) - y_i^-)^2 \quad (7)$$

$$\nabla_{\theta} J(\theta) \approx \frac{1}{N} \sum_i \nabla_a Q_\phi(S_i^C, a) \Big|_{S^C=S_i^C, a=\mu_\theta(S_i^A)} \cdot \nabla_{\theta} \mu_\theta(S_i^A) \Big|_{S^A=S_i^A} \quad (8)$$

액터-크리틱 구조의 심층 신경망을 학습시키기 용이하게 하기 위해 리플레이 버퍼(replay buffer)를 사용한다. 리플레이 버퍼는 학습 과정에서 매 단계(step)의 (상태 S_i^C , 상태 S_i^A , 행동 a_i , 비용 c_i , 다음 상태 S_{i+1}^C , 다음 상태 S_{i+1}^A) 튜플을 저장하는 유한한 크기의 캐시(cache)를 의미한다. off-policy 알고리즘이기 때문에 리플레이 버퍼를 크게 설정할 수 있다. 이는 데이터 간의 상관관계를 최소화하여 학습을 할 수 있도록 도와주는 효과를 가진다. 학습을 시작하기 전에 랜덤하게 행동 a 를 선택하는 정책을 기반으로 에피소드를 진행하여 초기 리플레이 버퍼를 생성시키는 과정을 진행한다.

심층강화학습 알고리즘으로 연속 행동 공간을 다루기 위해서는 탐험(exploration) 정책을 세우는 것이 필요하다. DDPG는 off-policy이기 때문에 탐험 문제를 학습 알고리즘과는 독립적으로 다룰 수 있기 때문에 간단한 방법으로 해결할 수 있다. 행동 $\mu_\theta(S)$ 에 노이즈 프로세스 N 에서 샘플링한 값을 더하여 행동하는 탐험 정책 $\mu'(S) = \mu_\theta(S) + N$ 을 사용할 수 있다. 본 연구에서는 N 을 정규분포 $N(0, \sigma_e^2)$ 를 따르도록 설정하였다. 에피소드가 점차 진행되면서 σ_e 가 감소하는 간단한 형태를 취해 학습의 효율을 높였다. 도출된 Algorithm 1의 pseudo-code는 제시된 것과 같다.

Algorithm 1. Inventory based Deep Deterministic Policy Gradient (IDDPG)

Randomly initialize critic network $Q_\phi(S^C, a)$ and actor network $\mu_\theta(S^A)$ with parameter ϕ and θ ;
 Initialize target network $Q_{\phi^-}(S^C, a)$ and $\mu_{\theta^-}(S^A)$ with $\phi^- \leftarrow \phi$ and $\theta^- \leftarrow \theta$;
 Initialize replay buffer R ;
 for episode = 1 : M do
 Initialize a random process N for exploration;
 for $t = 0 : T$ do
 Select action $a_t = \mu_\theta(S_t^A) + N_t$ according to the current policy and exploration noise;
 Execute action a_t and observe cost c_t and observe new state S_{t+1}^C and S_{t+1}^A ;
 Store transition $(S_t^C, S_t^A, a_t, c_t, S_{t+1}^C, S_{t+1}^A)$ in Buffer R ;
 Sample a random minibatch of N transitions $(S_i^C, S_i^A, a_i, c_i, S_{i+1}^C, S_{i+1}^A)$ from Buffer R ;
 Set target value $y_i^- = c_i + \gamma Q_{\phi^-}(S_{i+1}^C, \mu_{\theta^-}(S_{i+1}^A))$;
 Update the critic network by minimizing:

$$L(\phi) := \frac{1}{N} \sum_i (Q_\phi(S_i^C, a_i) - y_i^-)^2$$
;
 Update the actor network by using the sampled policy gradient:

$$\nabla_{\theta} J(\theta) \approx \frac{1}{N} \sum_i \nabla_a Q_\phi(S_i^C, a)|_{S^C=S_i^C, a=\mu_\theta(S_i^A)} \nabla_{\theta} \mu_\theta(S_i^A)|_{S^A=S_i^A}$$
;
 Update the target networks: $\theta^- \leftarrow \tau\theta + (1-\tau)\theta^-$ and $\phi^- \leftarrow \tau\phi + (1-\tau)\phi^-$;
 end for
end for

5. 수치 실험

본 장에서는 수치 실험을 통해 DDPG와 IDDPG의 성능을 비교한다. 우선 최적해가 알려져 있는 재고관리 문제에 일반적인 수요분포 하에서의 리드타임이 일정한 무한

기간 뉴스벤더 문제의 실험을 통해 DDPG와 IDDPG를 비교한다. 이때 재고관리 문제의 복잡도를 높이는 요인으로 작용하는 리드타임, 수요의 변동성, 그리고 재고부족비용을 변경하면서 실험 환경을 구성하고 각 요인에 변화에 따른 민감도를 분석한다. 또한, 최적해가 알려져 있지 않은 비정형 수요(독립항등분포 가정과 지수족 가정이 성립하지 않는 복잡한 수요)를 가진 재고관리 문제에 IDDPG를 적용하고 SAA와 IDDPG를 비교함으로써, 최적해가 알려져 있지 않은 재고관리 문제에 대해서도 IDDPG를 확장하여 적용할 수 있는지 그 가능성을 확인해보고자 한다.

5.1 실험 설계

심층강화학습을 활용해 뉴스벤더 문제를 풀기 위해서는 두 가지 과정이 필요하다. 마르코프 결정 과정 모델을 반영한 실험 환경(environment)을 설계 및 구현하는 과정을 거치고, 문제에 적합한 심층강화학습의 초매개변수(hyperparameter) 설정값을 알기 위해 튜닝 과정(tuning process)을 진행해야 한다. 과정을 진행한다.

상태 (S_t^C, S_t^A) 와 행동 a_t 에 대해서 비용 c_t 와 다음 상태 (S_{t+1}^C, S_{t+1}^A) 가 뉴스벤더 문제 상황에 따라서 결정되도록 실험 환경을 구축하기 위해 수요 분포, 리드타임, 상태 공간, 행동 공간, 비용 함수 등을 결정해야 한다. 가장 기본이 되는 실험 환경을 설정하고 리드타임 l , 수요의 변동성 σ , 재고부족비용 p 를 각각 변화시키며 다양한 환경에서의 실험을 진행하고자 한다. 기본이 되는 실험 환경 조건은 $l = 4, \sigma = 1.0, p = 4$ 로 설정했다. 리드타임 $l \in \{0, 2, 4, 6, 8\}$, 수요의 변동성 $\sigma \in \{0.5, 1.0, 1.5, 2.0, 2.5\}$, 재고부족비용 $p \in \{1, 4, 9, 16, 25\}$ 에 대한 민감도 분석 실험을 <Table 1>과 같이 진행한다. 따라서, 기본 조건의 중복을 제외하면 13가지의 실험 환경을 다루었다.

수요 분포는 정규분포 $N(\mu, \sigma^2)$ 를 $[0, \infty)$ 구간으로 자른 잘린정규분포(truncated normal distribution)를 따른다고 가정했다. 상태 공간과 행동 공간은 수요 분포와 리드타임 l 에 따라서 적절하게 설정하는 것이 중요하다. 행동 공간은 $a_l = -\mu$ 와 $a_u = (l+1)\mu + (3l+3)\sigma$ 로 설정하여 $a \in [a_l, a_u]$ 를 따르도록 하였다. 상태 공간은 이에 맞추어 $s_0^C \in [-a_u, (\frac{1}{2}l+1)a_u]$, $s_1^C, \dots, s_{l-1}^C \in [0, a_u]$, $S^A \in [-a_u, \frac{3}{2}la_u]$ 를 따르도록 하였다. 재고유지비용 h 를 1로 고정하고 재고부족비용 p 를 변화시키며 비용 함수를 설정하였다.

Table 1. Sensitivity experiment conditions

Sensitivity Factor	Lead Time (l)	Demand Volatility (σ)	Shortage Cost (p)
Lead Time	{0, 2, 4, 6, 8}	1.0	4
Demand Volatility	4	{0.5, 1.0, 1.5, 2.0, 2.5}	4
Shortage Cost	4	1.0	{1, 4, 9, 16, 25}

DDPG와 IDDPG에는 심층 신경망 층의 수, 층별 뉴런의 수, 층별 활성화 함수(activation function), 액터 네트워크의 학습률(learning rate), 크리틱 네트워크의 학습률, 타겟 네트워크 파라미터 τ 등의 초매개변수가 존재한다. 선행 연구들과 실험 경험을 토대로 초매개변수들의 튜닝 실험 여부와 범위를 설정하였다. 각 신경망 층에서 ReLU를 활성화 함수로 사용하며 [150, 120, 80, 20]구조의 4개의 신경망 층을 가지는 네트워크를 고정적으로 사용하기로 결정했다. 액터 네트워크의 학습률, 크리틱 네트워크의 학습률, 타겟 네트워크 파라미터 τ 는 아래의 <Table 2>와 같이 각각 7, 7, 3가지의 값을 가지도록 하여 147가지 조합을 실험하여 튜닝하였다.

실험 환경마다 적합한 초매개변수 설정값이 다르므로 147가지 조합의 초매개변수에 해당하는 DDPG와 IDDPG 알고리즘을 13가지 실험 환경에 맞추어 학습시켜 성능을 비교한다. 즉, 튜닝 과정에서 다루어야 하는 실험 환경과 알고리즘의 조합은 3,822가지가 된다. 각 조합에 대해서 100,000 스텝에 해당하는 시뮬레이션을 진행한다. 각 학습 과정의 초기에는 행동 공간의 랜덤 샘플링 정책을 통해 10,000 스텝을 진행하여 리플레이 버퍼를 초기화한다. 그리고 학습 에피소드의 최대 길이를 500 스텝으로 제한하며, 알고리즘의 평가를 1,000 스텝마다 진행한다. 알고리즘의 평가 과정에서는 100개의 평가 에피소드를 진행하며 발생된 스텝의 평균 비용을 계산한다. 모든 에피소드의 초기 상태는 $s_0^C \in [-a_u + \max(1, l)\mu, a_u - \max(1, l)\mu]$, $s_1^C, \dots, s_{l-1}^C \in [0, 2\mu]$ 에서의 랜덤 샘플링을 통해서 설정하

였다.

심층강화학습의 특성상 스텝이 진행됨에 따라서 평가 과정에서의 평균 비용 값이 낮아지지만 하는 것이 아니라 늘어나는 경우도 존재한다. 중간 스텝에 학습이 잘 진행되어 평균 비용 값이 굉장히 낮게 나오다가도 탐험 정책에 의해 학습의 방향이 잘못되어 후반 스텝에서는 평균 비용 값이 높게 나오는 문제가 발생될 수 있기 때문이다. 이런 경우들이 존재하기 때문에 단순히 마지막 스텝에서의 네트워크를 학습결과로 바라보기 힘들다. 딥러닝의 경우 중간에 성능이 좋더라도 후반에 나빠지는 경우 과적합(overfitting)에 대한 우려로 인해서 중간에 성능이 좋았을 때의 네트워크를 학습결과로 바라보기 힘들 수 있다. 그러나, 심층강화학습의 특성상 학습 데이터와 실제 상황과의 괴리가 존재하지 않기 때문에 과적합에 대한 우려를 가지지 않아도 된다. 따라서, 심층강화학습의 튜닝 과정에서의 평가 지표로 평가 스텝 중에서 낮은 평균 비용을 보이는 스텝을 가지는 조합이 좋다고 판단하는 것이 합리적이다. 본 연구에서는 100,000개의 스텝에서 진행되는 100개의 평가 과정에서 가장 좋은 평균 비용을 제시한 10개의 평가 과정의 평균을 평가 지표로 사용하여 튜닝을 하였다.

모든 조합에 대한 시뮬레이션을 진행하여 각 실험 환경마다 DDPG와 IDDPG에 대해 평가 지표가 좋은 상위 3개의 초매개변수 설정 값을 선정했다. 선정된 조합들에 대해서 3번의 시뮬레이션을 추가적으로 진행하여 최종적으로 가장 좋은 설정 값을 선정하는 방식으로 튜닝을

Table 2. Hyperparameters of the DDPG and IDDPG that require tuning

Hyperparameters	Range for tuning
Number of layers	4
Width of layers	[150, 120, 80, 20]
Activation functions	ReLU
Actor Learning rate (10^{-x})	$x \in \{3.0, 3.5, 4.0, 4.5, 5.0, 5.5, 6.0\}$
Critic Learning rate (10^{-x})	$x \in \{2.0, 2.5, 3.0, 3.5, 4.0, 4.5, 5.0\}$
Tau (10^{-x})	$x \in \{2.0, 3.0, 4.0\}$

진행했다. 튜닝이 완료된 조합들에 대해서 총 10번의 시뮬레이션을 진행하여 평가 지표의 평균을 최종적인 알고리즘 성능으로 평가하였다. 이 때, 시뮬레이션 과정에서 네트워크 초기 설정값과 탐험 정책으로 인해 학습에 문제가 발생한 경우에 대해서는 새롭게 시뮬레이션을 진행하였다.

5.2 뉴스벤더 문제에서의 DDPG와 IDDPG의 성능 비교

튜닝 과정을 통해서 각 실험 환경과 알고리즘에 적합한 초매개변수 조합은 <Table 3>과 같다. 평균적으로 DDPG는 10^{-4} 의 액터 네트워크 학습률, $10^{-2.5}$ 의 크리틱 네트워크 학습률, $10^{-3.0}$ 의 타겟 네트워크 파라미터 τ 를 기준으로 약간의 변화를 주는 조합이 좋은 성능을 보였다. IDDPG는 $10^{-5.5}$ 의 액터 네트워크 학습률, $10^{-3.0}$ 의 크리틱 네트워크 학습률, $10^{-2.0}$ 의 타겟 네트워크 파라미터 τ 를 기준으로 약간의 변화를 주는 조합이 좋은 성능을 보였다. 특히 액터 네트워크 학습률에서 큰 차이를 보여주는 것을 알 수 있는데 이는 액터 네트워크의 상태 변수를 다르게 정의한 점에서 비롯된 것으로 보인다. 상태 변수의 차원을 낮춤으로써 학습률이 낮아 전역 최적해가 아닌 지역 최적해에서 빠져나오지 못하는 문제가 발생할 확률이 낮아졌기 때문이라고 생각된다.

<Table 3>에 해당하는 초매개변수 조합을 기준으로 진행된 실험 결과는 <Table 4>와 같다. 리드타임이 0인 경우를 제외한 모든 실험 환경에서 IDDPG가 DDPG보다 최적해에 가까운 해를 제시하며 좋은 성능을 보였다. 리드타임이 증가할수록, 수요의 변동성이 커질수록, 재고유지비용 대비 재고부족비용의 비율이 증가할수록 IDDPG의 성능 개선폭이 증가하는 경향을 보인다. 액터 네트워크의 상태 변수 차원 감소 폭이 리드타임이 증가할수록 커지기 때문이다. 주문량 0에 대응되는 행동 공간을 확장시킨 효과가 수요의 변동성과 함께 커지기 때문으로 해석된다. 상태공간을 벗어나는 경우 패널티를 추가하는 방향으로 비용 함수를 변형시킨 것의 효과가 재고부족비용과 함께 커지는 것으로 보인다.

모든 상태에 대한 최적해를 구할 수 있는 뉴스벤더 문제 상황에서 실험을 진행했기 때문에 단순히 평균 비용뿐만 아니라 상태별 행동 패턴에 대한 분석이 가능하다. 민감도 분석의 기준인 리드 타임, 수요의 변동성, 재고부족비용에 따라서 각각 <Fig. 1>, <Fig. 2>, <Fig. 3>으로 실험 결과를 나누어 정리했다. 각 그림의 (a)는 성능 곡선 그래프를, (b)는 상태-행동 밀도 곡선 그래프를 나타낸다.

성능 곡선 그래프는 10번의 시뮬레이션 과정에서 학습 스텝이 진행됨에 따라 평균 비용이 어떻게 변화되는

Table 3. Well-performing hyperparameter set

Environment			DDPG			IDDPG		
Lead Time	Demand Volatility	Shortage Cost	Actor Learning rate	Critic Learning rate	Tau	Actor Learning rate	Critic Learning rate	Tau
0	1.0	4	4.0	4.0	4.0	5.5	3.5	4.0
2	1.0	4	3.5	3.0	4.0	6.0	3.0	2.0
4	1.0	4	3.5	2.0	3.0	5.0	3.0	2.0
6	1.0	4	4.0	2.5	3.0	5.5	3.0	2.0
8	1.0	4	4.0	3.0	3.0	5.0	2.0	2.0
4	0.5	4	4.0	2.0	3.0	4.0	3.0	2.0
4	1.0	4	3.5	2.0	3.0	5.0	3.0	2.0
4	1.5	4	3.5	3.0	3.0	4.5	3.5	2.0
4	2.0	4	3.5	3.0	4.0	5.5	2.5	2.0
4	2.5	4	4.5	2.5	4.0	5.5	3.0	2.0
4	1.0	1	4.0	3.0	3.0	5.0	3.0	2.0
4	1.0	4	3.5	2.0	3.0	5.0	3.0	2.0
4	1.0	9	3.5	3.0	4.0	5.5	3.0	2.0
4	1.0	16	5.0	2.5	3.0	5.5	3.0	2.0
4	1.0	25	4.5	3.5	4.0	5.5	2.5	2.0

지를 의미한다. 진한 곡선은 시물레이션들의 평균을, 연한 영역은 시물레이션들의 평균을 기준으로 표준편차를 가감한 범위를, 상수 함수는 최적해의 평균 비용값을 의미한다. 상태-행동 밀도 곡선 그래프는 10번의 시물레이션마다 평균 비용이 가장 낮았던 3번의 평가 스텝의 네트워크를 사용하여 다양한 상태에서 어떤 행동을 하는지 확인한 것이다. 30가지 액터와 크리틱 네트워크 쌍마다 10,000번의 상태 샘플에서 어떤 행동을 선택하는지를 시물레이션 해보았다. 각 실험 환경마다 300,000개의 상태-행동 데이터를 선형 그래프로 정리하였다. 성능 곡선 그래프와 마찬가지로 진한 곡선은 각 상태별 선택한 행동들의 평균이며 연한 영역은 평균에 표준편차를 가감한 범위를 의미한다. 그리고 각진 선은 상태별 최적의 행동값을 의미한다.

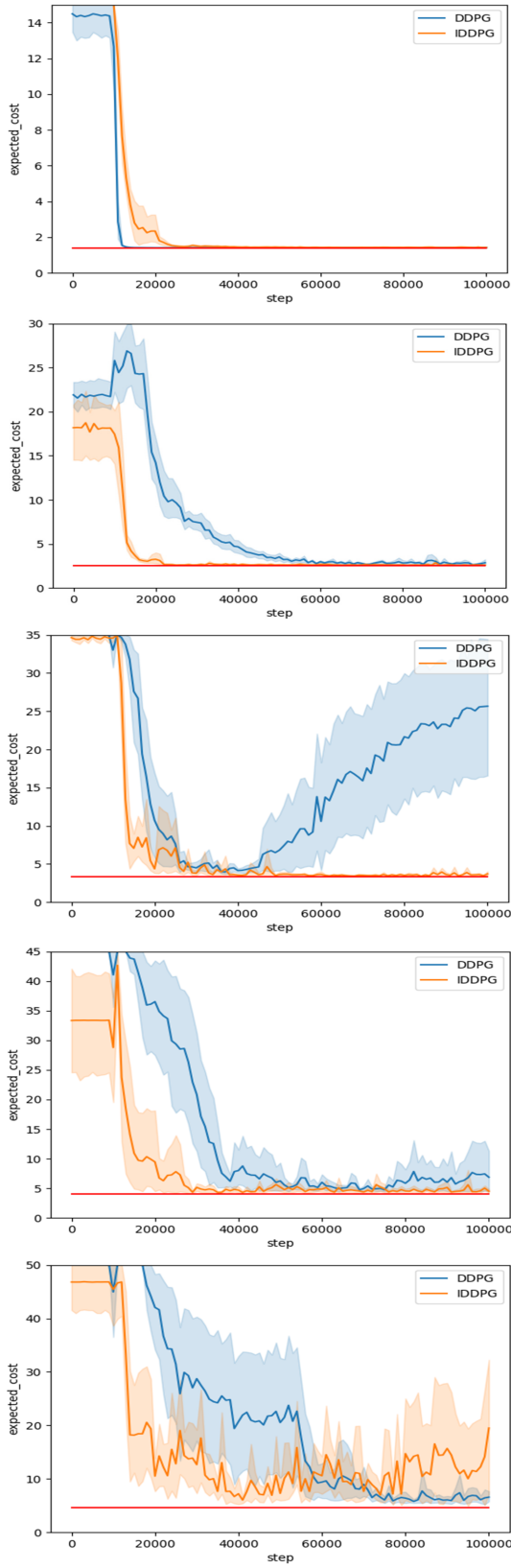
성능 곡선 그래프를 보면 대부분의 경우 IDDPG가 DDPG보다 학습 속도가 빠르며, 안정적으로 최적해에 가까운 성능을 보이는 것을 볼 수 있다. 상태-행동 밀도 곡선을 보면 DDPG의 경우 같은 상태에 놓여졌을 때에도 굉장히 다양한 행동을 선택하는 것을 볼 수 있다. 이는 액터 네트워크의 상태 변수가 기말재고와 운송중 재고를 모두 입력값으로 활용하기 때문에 재고상태가 같아도 다양한 행동을 선택하는 경향을 보이기 때문으로 해석된

다. 이에 반해 IDDPG는 변동성이 낮은 정책을 제시하는 것을 볼 수 있다. 상태 공간의 경계에서 DDPG가 제시하는 행동 값이 최적해와 크게 차이나는 모습을 볼 수 있다. 반면, IDDPG는 경계에서도 최적해에 비슷한 행동 값을 선택하는 모습을 볼 수 있다. 다만, IDDPG의 정책을 보면 상대적으로 낮은 상태 범위에서 최적해보다 높은 행동을 선택하는 경향성을 보이는데 이는 상태공간을 벗어나면 추가되는 패널티의 값을 크게 잡았기 때문으로 보인다.

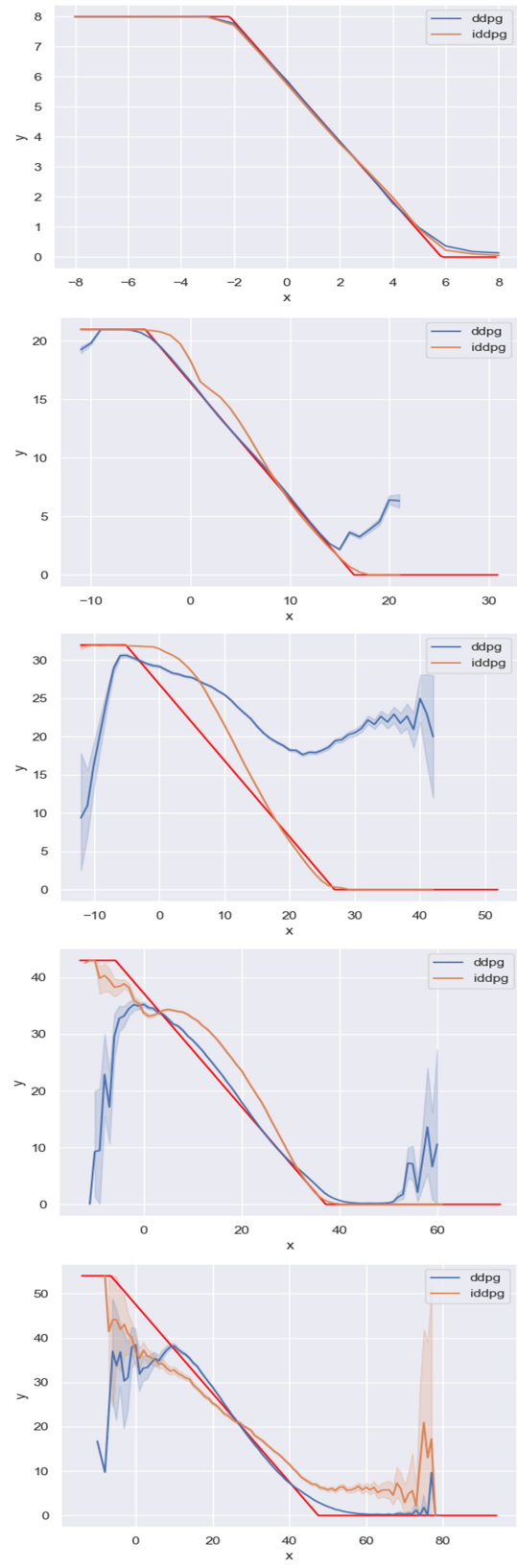
DDPG와 IDDPG 모두 최적해와 완전 동일한 형태의 정책을 제시하고 있지는 못함에도 불구하고 최적해에 가까운 성능을 보이는 이유는 에피소드의 주기가 지남에 따라 최적해를 제시하는 상태공간에서 스텝이 진행되기 때문이다. 이러한 현상은 심층강화학습 알고리즘의 장점이자 단점이 될 수 있다고 보여진다. 빈번하게 마주하게 되는 상태에서도 최적해를 제시할 수 있다는 점은 장점으로 생각된다. 그러나, 모든 상태에서 최적의 행동을 선택할 수 있는 방향으로 추가적인 연구가 진행되어야 한다. 뿐만 아니라, 최적해를 모르는 문제에 대해서 알고리즘을 평가할 때 단순히 평균 비용을 지표로 사용하여 성능을 평가하는 것만으로는 부족할 수 있음을 시사한다. 따라서, 본 연구처럼 최적해를 아는 상태에서 알고리즘의 개선점을 엄밀하게 보이는 과정이 필요하다.

Table 4. Experiment results

Environment			Basestock	DDPG		IDDPG	
Lead Time	Demand Volatility	Shortage Cost	Optimal Expected Cost	Average Expected Cost	Gap (%) compared to Basestock	Average Expected Cost	Gap (%) compared to Basestock
0	1.0	4	1.398	1.399	0.1	1.409	0.8
2	1.0	4	2.513	2.596	3.3	2.529	0.6
4	1.0	4	3.320	3.644	9.8	3.334	0.4
6	1.0	4	3.994	4.309	7.9	4.041	1.2
8	1.0	4	4.619	5.243	13.5	4.733	2.5
4	0.5	4	1.741	1.854	6.5	1.759	1.0
4	1.0	4	3.320	3.644	9.8	3.334	0.4
4	1.5	4	4.898	5.770	17.8	4.954	1.1
4	2.0	4	6.374	7.134	11.9	6.410	0.6
4	2.5	4	7.715	9.050	17.3	7.758	0.5
4	1.0	1	1.833	2.105	14.9	1.849	0.9
4	1.0	4	3.325	3.644	9.8	3.334	0.4
4	1.0	9	4.356	4.836	11.0	4.380	0.5
4	1.0	16	5.256	25.316	381.7	5.285	0.5
4	1.0	25	6.101	9.212	51.0	6.102	0.0



(a) Performance curves



(b) Density plots

Fig. 1 Performance curves and Density plots for $l \in \{0, 2, 4, 6, 8\}$, $\sigma = 1.0$, $p = 4$

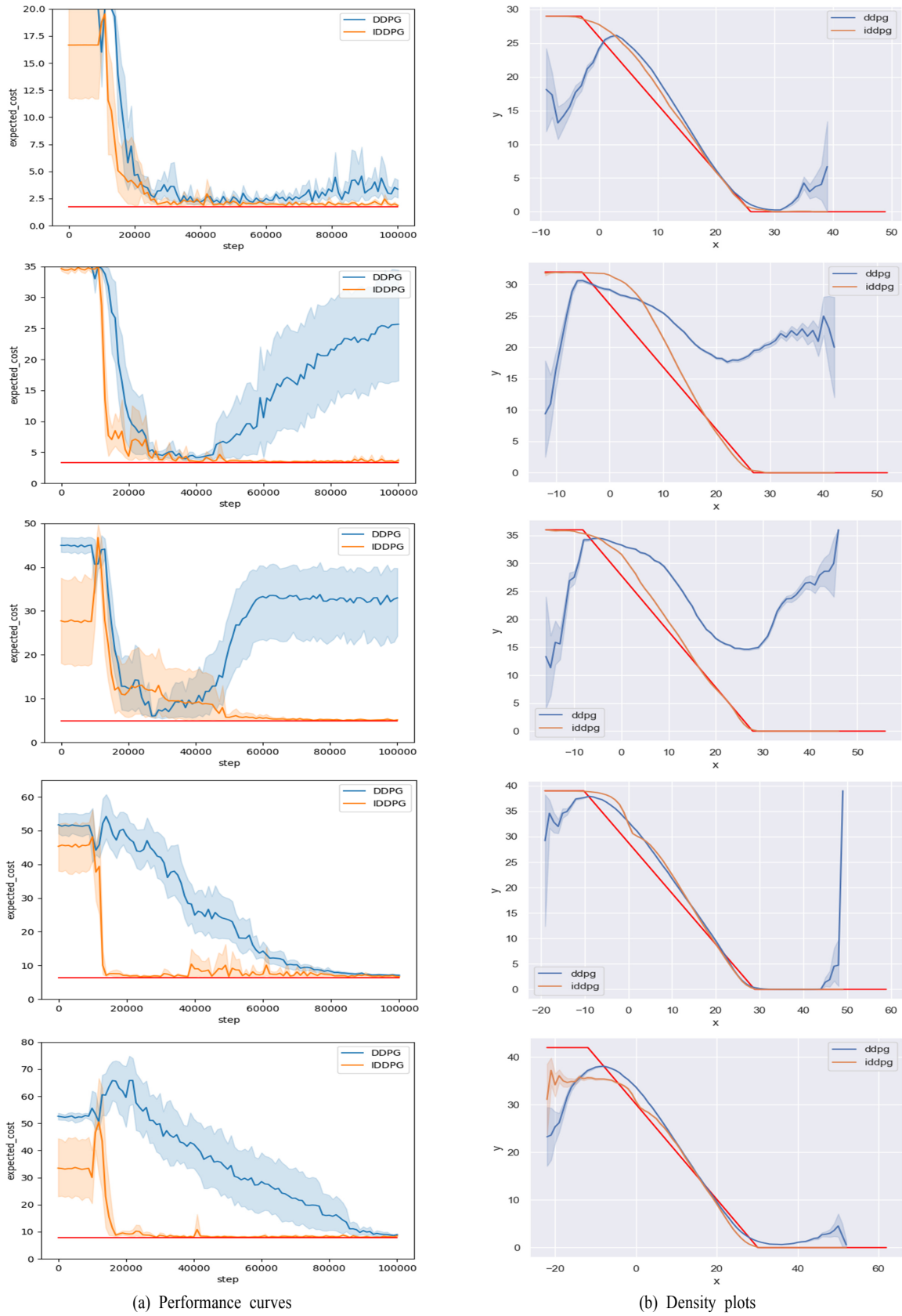
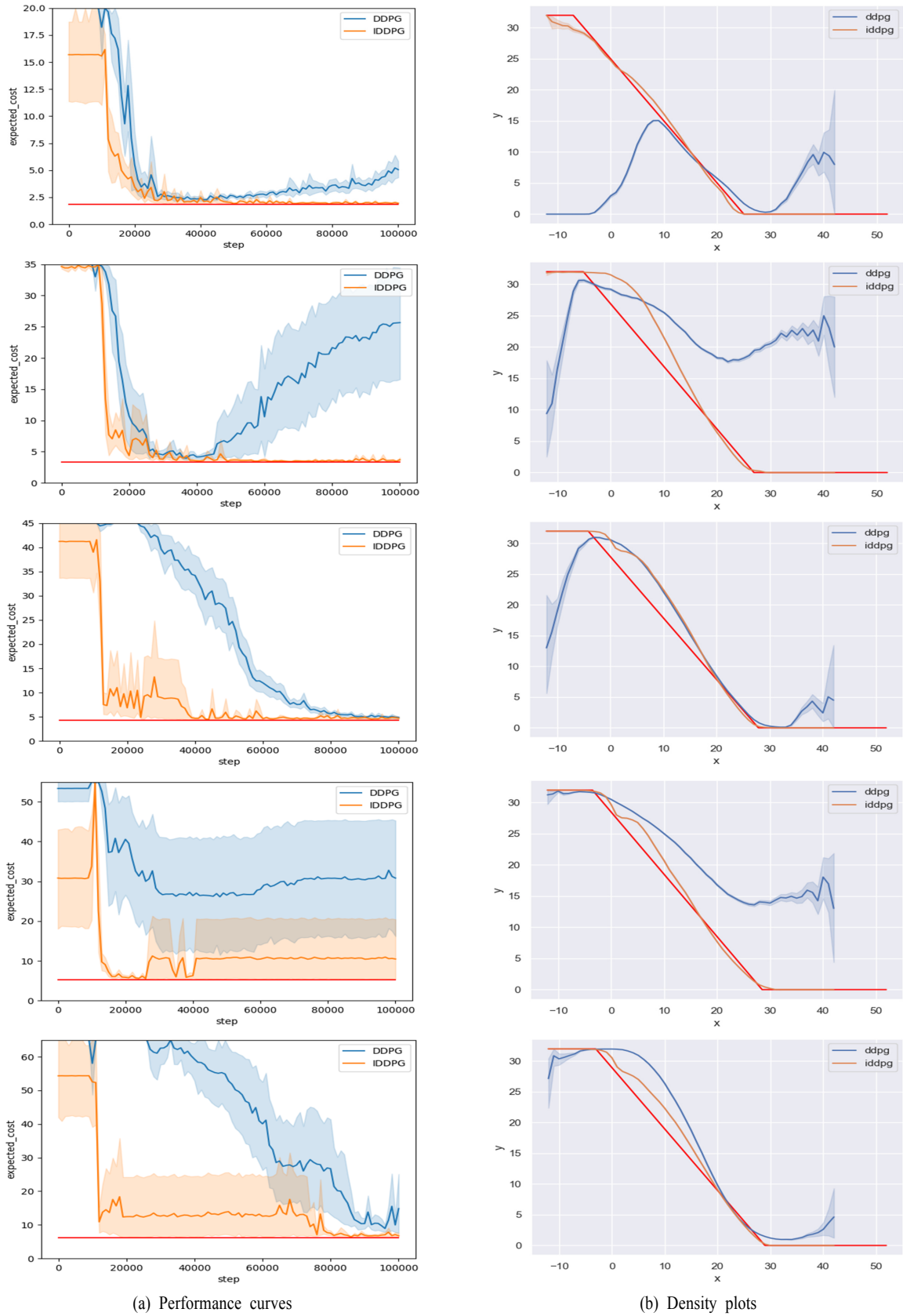


Fig. 2 Performance curves and Density plots for $l=4$, $\sigma \in \{0.5, 1.0, 1.5, 2.0, 2.5\}$, $p=4$



5.3 비정형 수요의 뉴스벤더 문제에서의 DDPG와 IDDPG의 비교

독립항등분포 성질을 만족하는 수요 분포를 가정하면 주문 고정 비용과 리드타임이 존재하는 뉴스벤더 문제에서 (s, S) 정책이 최적 정책이다. 그러나, 흥미롭게도 (s, S) 정책을 통해서 발생하는 주문 데이터는 독립항등분포 성질을 만족하지 못한다. (s, S) 정책을 사용하는 고객을 대상으로 재고관리 의사결정을 해야 하는 입장에서는 독립항등분포와 지수족 가정이 만족되지 않는 비정형 수요 분포를 다루어야 한다. (s, S) 정책에 의한 주문 데이터의 흥미로운 점은 주문량이 0인 경우가 빈번하게 발생한다는 점이다. 실제로 현실의 주문 데이터를 보면 주문 고정 비용과 최소 주문량과 같은 제약 때문에 한번에 많은 양의 주문을 하는 경우가 많다.

본 연구에서는 이러한 상황을 반영한 실험 환경에서 SAA와 IDDPG의 성능을 비교해보고자 한다. 고객의 수요는 $N(5, 2^2)$ 를 $[0, \infty)$ 구간으로 자른 잘린정규분포를 따르며 $(20, 40)$ 의 정책을 통해 주문한다고 가정하였다. SAA는 리드타임이 0인 경우에만 적용가능 하므로 실험 환경은 리드타임 $l=0$ 으로 설정하였다. 또한, 재고유지 비용 $h=0$, 재고부족비용 $p=4$ 으로 설정하였다.

IDDPG가 100,000 스텝을 통해 학습을 하기 때문에 SAA도 100,000개의 샘플 데이터를 가지고 있다고 가정하였다. 이 때, SAA는 항상 재고상태를 20으로 유지하는 정책을 사용한다.

IDDPG는 SAA와 달리 다양한 정보를 추가적으로 활용하기 용이하다는 장점이 있다. 본 연구에서는 고객의 최근에 주문량이 0인 기간의 길이 u 정보를 활용할 수 있도록 기존의 상태 변수를 변형했다. 이항변수 벡터 u_t 는 주기 t 에서 고객의 최근에 주문량이 0인 기간의 길이에 해당하는 원소는 1, 아닌 원소는 0을 가지는 벡터다. 문제의 편의를 위해 $\bar{u}$ 를 충분히 큰 값으로 설정하고, 고객이 주문을 미룰 수 있는 최대 기간을 $\bar{u}$ 로 설정하였다. 즉, $\bar{u}$ 기간동안 주문을 안 한 경우에는 재고상태가 20보다 큰 경우에도 40이 되도록 주문을 하도록 하였다

앞서 정의한 대로, 기존의 상태 변수 $S_t = (s_0 = I_{t-1} + q_{t-l}, s_1 = q_{t-l+1}, s_2 = q_{t-l+2}, \dots, s_{l-1} = q_{t-1})$ 의 뒤에 $(\bar{u}+1)$ 차원 이항변수 벡터 $u_t = (u_{t0}, u_{t1}, \dots, u_{t\bar{u}})$ 를 추가하는 형태인 $S_t = (s_0 = I_{t-1} + q_{t-l}, s_1 = q_{t-l+1}, s_2 = q_{t-l+2}, \dots, s_{l-1} = q_{t-1}, s_l = u_{t0}, s_{l+1} = u_{t1}, \dots, s_{l+\bar{u}} = u_{t\bar{u}})$ 으로 확장한다.

실험 예제에는 $10^{-5.5}$ 의 액터 네트워크의 학습률, $10^{-3.5}$ 의 크리틱 네트워크의 학습률, $10^{-3.0}$ 의 타겟 네트워크 파라미터 τ 초매개변수 조합으로 구성된 IDDPG를 사용하여 3번의 시뮬레이션을 진행했다. 실험 결과를 성능 곡선 그래프로 나타내면 <Fig. 4>와 같다. 상수 함수 선은 SAA의 평균 비용을, 진한 곡선은 시뮬레이션들의 평균을, 연한 영역은 평균을 기준으로 표준편차를 가감한 범위이다. 그래프에서 볼 수 있듯이 IDDPG가 압도적으로 좋은 성능을 보인다는 것을 알 수 있다. IDDPG는 SAA의 평균 비용 18.940과 비교하여 68.3% 낮은 6.012의 평균 비용을 보인다.

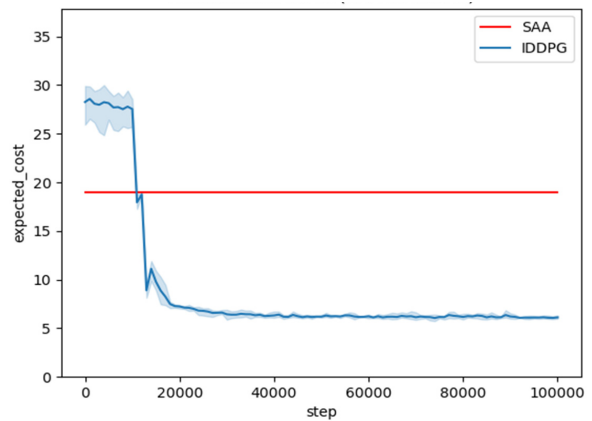


Fig. 4 Performance curve (SAA vs. IDDPG)

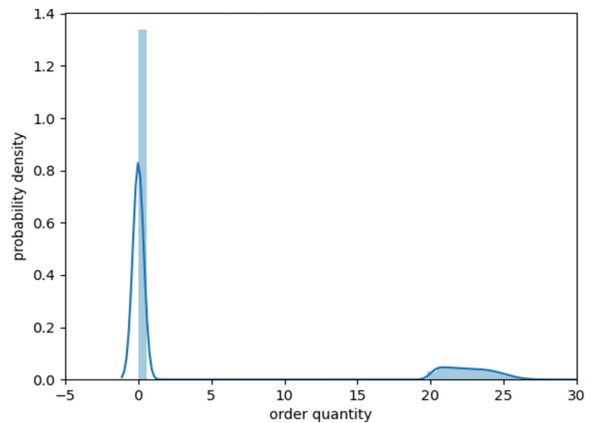


Fig. 5 (s, S) order distributions

SAA와 IDDPG 사이의 성능 차이는 (s, S) 정책에 의한 고객의 주문 데이터 분포를 보면 이해할 수 있다. 100,000개의 수요 데이터를 히스토그램으로 그려보면 <Fig. 5>와 같다. 이러한 수요 데이터를 정보 u 의 값에 따라 나누어 히스토그램을 그려보면 <Fig. 6>과 같다. 정보 u 의 값에 따라 수요의 분포가 크게 달라지는 것을 볼 수

있다. 그러나, SAA는 수요분포의 독립항등분포성에 대한 가정이 필수적이기 때문에 정보 u 의 값에 따라 달라지는 분포의 특성을 고려하기 힘들다. 따라서, SAA는 주문량이 0일 확률이 굉장히 높은 $u=0, 1, 2$ 의 경우에도 $u=3, 4, 5$ 일 때와 마찬가지로 재고상태를 20으로 유지하기 위해 주문하는 비효율적인 모습을 보여준다. 반면에, IDDPG는 정보 u 의 값에 따라 다르게 주문하는 최적의 의사결정 정책을 학습하기 때문에 좋은 성능을 보인다.

것으로 보인다.

엄밀한 실험을 진행하지는 못했지만, 간단한 실험 예제를 통해서 IDDPG의 확장성이 좋기 때문에 다양한 정보를 쉽게 추가할 수 있다는 점을 보였다. 뿐만 아니라 기존의 데이터 기반 의사결정에 비해 좋은 성능을 보일 수 있음을 보였다. 이로 인해서 최적해를 구하기 어려운 재고관리 문제로의 확장 가능성을 제시하였다는 데에 연구의 의의를 둘 수 있다.

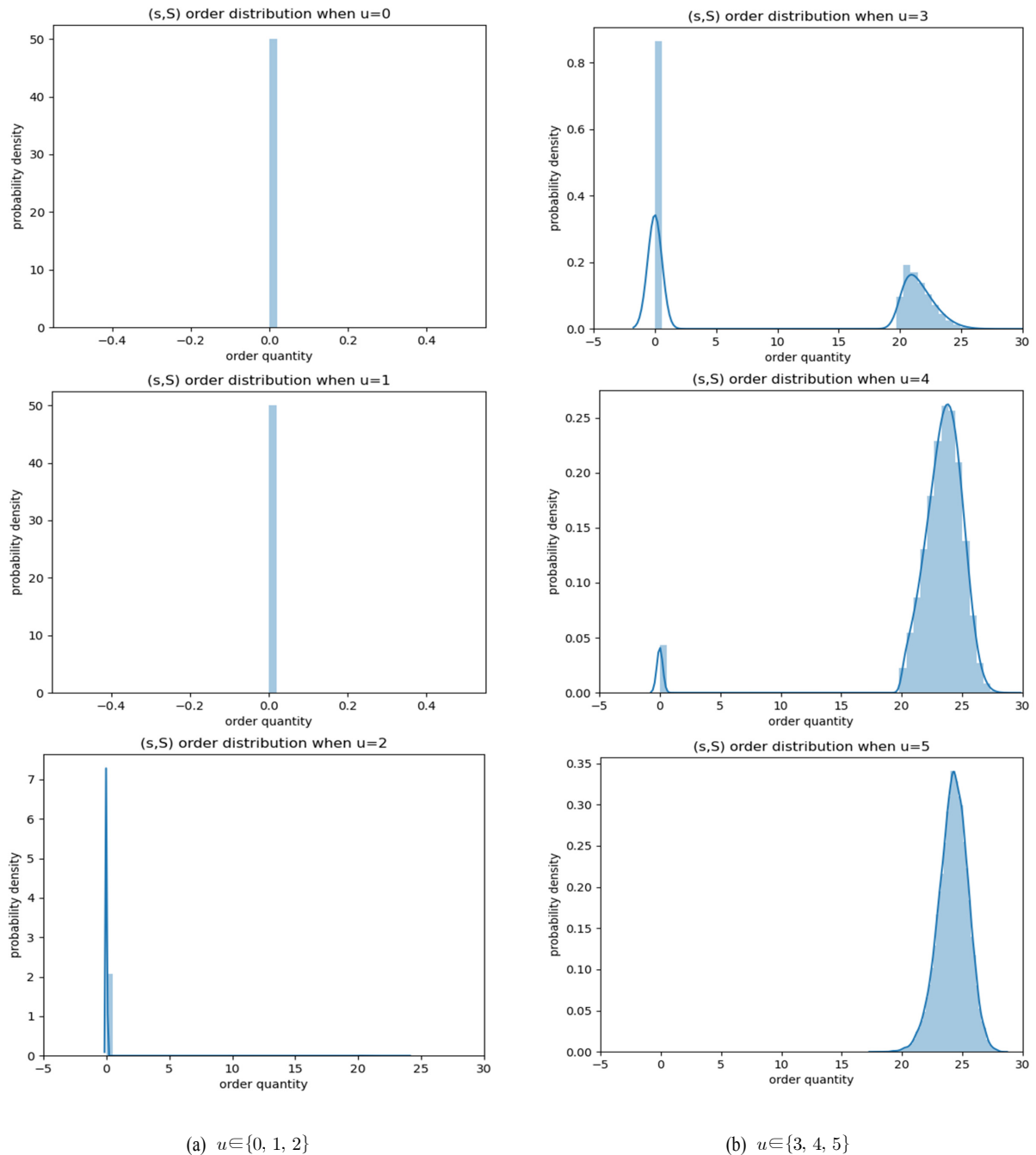


Fig. 6. (s, S) order distributions for specified $u \in \{0, 1, 2, 3, 4, 5\}$

6. 결 론

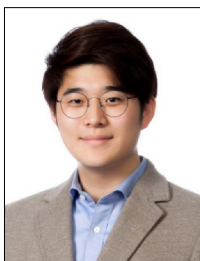
본 연구에서는 재고관리 문제의 특성을 반영한 심층 강화학습 알고리즘으로 Inventory-based DDPG(IDDPG)를 제안하였다. 본 알고리즘은 연속 행동 공간을 다루는 대표적인 심층강화학습 알고리즘인 DDPG를 재고관리 문제의 특성들을 고려하여 상태 변수, 행동 공간, 비용 함수 등을 적절히 변형시켰다. 본 알고리즘의 성능을 분석하기 위해 기존에 최적해가 알려져 있지만 복잡한 재고관리 문제인 리드타임이 일정한 무한 기간 뉴스벤더 문제(infinite horizon newsvendor problem with constant lead time)에 수치실험을 수행하여 DDPG와 IDDPG가 도출하는 해의 성능을 비교하였고, IDDPG가 DDPG 대비 최적해에 충분히 가까운 좋은 해를 도출함을 입증하였다. 또한 본 연구에서 제안된 IDDPG는 심층강화학습의 특징에 따라 기존의 알고리즘들이 확장되기 어려운 비정형 수요(수요의 독립성이나 지수족 가능성이 성립하지 않는 복잡한 수요 형태)에도 확장 가능성을 수치실험 예제를 통해 보여주었다.

본 연구는 심층강화학습 알고리즘이 재고관리 문제의 특성을 반영하여 조정된다면, 최적해에 가까운 좋은 성과를 기대할 수 있음을 보여주고 있다. 또한 기존의 연구가 다루기 어려운 비정형적 수요에도 확장 가능성을 보여줘 새로운 연구 방향을 제시한 의의가 있다. 본 연구의 결과를 바탕으로 추후 각종 비정형 수요 형태 하에서 IDDPG가 기존의 SAA등의 데이터기반 재고관리 알고리즘에 비하여 어느 정도 개선된 해를 도출할 수 있는지를 연구할 수 있을 것으로 기대한다. 또한 수요의 특성에 따라 IDDPG의 성능을 측정하고, 알고리즘을 개선하는 연구도 수행될 수 있을 것으로 생각한다.

REFERENCES

- [1] Arrow, K. J., Harris, T., & Marshak, J. (1951). Optimal Inventory Policy. *Econometrica*, 19(3), 250-272.
- [2] Arrow K. J., Karlin, S., & Scarf, H. (1958). *Studies in the Mathematical Theory of Inventory and Production*. Stanford University Press. Stanford, CA.
- [3] Atkinson, A. A. (1979). Incentives, Uncertainty, and Risk in the Newsboy Problem. *Decision Sciences*, 10(3), 341-357.
- [4] Ban, G. Y., & Rudin, C. (2019). The Big Data Newsvendor: Practical Insights from Machine Learning. *Operations Research*, 67(1), 90-108.
- [5] Bellman, R. (1954). The Theory of Dynamic Programming. *Bulletin of the American Mathematical Society*, 60(6), 503-515.
- [6] De Moor, B. J., Gijsbrechts, J., & Robert B. (2021). Reward Shaping to Improve the Performance of Deep Reinforcement Learning in Perishable Inventory Management. *European Journal of Operational Research*, forthcoming.
- [7] Edgeworth, F. Y. (1888). The Mathematical Theory of Banking. *Journal of the Royal Statistical Society*, 51(1), 113-127.
- [8] Fan, J., Wang, Z., Xie, Y., & Yang, Z. (2019). A Theoretical Analysis of Deep Q-Learning. Available at arXiv: 1901.00137.
- [9] Gijsbrechts, J., Boute, R. N., Van Mieghem, J. A., & Zhang, D. J. (2020). Can Deep Reinforcement Learning Improve Inventory Management? Performance on Dual Sourcing, Lost Sales and Multi-Echelon Problems. *Manufacturing & Service Operations Management*, Forthcoming. Available at SSRN: 10.2139/ssrn.3302881.
- [10] Haarnoja, T., Zhou, A., Abbeel, P., & Levine, S. (2018). Soft Actor-Critic: Off-Policy Maximum Entropy Deep Reinforcement Learning with a Stochastic Actor. Available at arXiv: 1801.01290.
- [11] Homem-De-Mello, T. (2000). Monte Carlo methods for discrete stochastic optimization. In: Uryasev, S., & Pardalos, P. M. (eds), *Stochastic Optimization: Algorithms and Applications*, 97-119, Springer, Boston, MA.
- [12] Kleywegt, A. J., Shapiro, A., & Homem-de-Mello, T. (2001). The Sample Average Approximation Method for Stochastic Discrete Optimization. *SIAM Journal on Optimization*. 12(2), 479-502.
- [13] Kumar, A., Fu, J., Tucker, G., & Levine, S., (2019). Stabilizing Off-Policy Q-Learning via Bootstrapping Error Reduction. *Advances in Neural Information Processing Systems 2019*, Available at arXiv: 1906.00949.
- [14] Lau, H. S. (1980). The Newsboy Problem Under Alternative Optimization Objectives. *Journal of the Operational Research Society*, 31(6), 525-535.
- [15] Lau, A. H. L. & Lau, H. S. (1988). The Newsboy Problem with Price-Dependent Demand Distribution. *IIE Transactions*, 20(2), 168-175.
- [16] Levi, R., Roundy, R. O., & Shmoys, D. B. (2007). Provably

- Near-Optimal Sampling-Based Policies for Stochastic Inventory Control Models. *Mathematics of Operations Research*, 32(4), 821-839.
- [17] Levi, R., Perakis, G., & Uichanco, J. (2015). The Data-Driven Newsvendor Problem: New Bounds and Insights. *Operations Research*, 63(6), 1294-1306.
- [18] Li, Y. (2018). Deep Reinforcement Learning. Available at arXiv: 1810.06339.
- [19] Lillicrap, T. P., Hunt, J. J., Pritzel, A., Heess, N., Erez, T., Tassa, Y., Silver, D., & Wierstra, D. (2015). Continuous Control with Deep Reinforcement Learning. Available at arXiv: 1509.02971.
- [20] Liyanage, L. H. & Shanthikumar, J. G. (2005). A Practical Inventory Control Policy Using Operational Statistics. *Operations Research Letters*, 33(4), 341-348.
- [21] Mnih, V., Kavukcuoglu, K., Silver, D., Graves, A., Antonoglou, I., Wierstra, D., & Riedmiller, M. (2013). Playing Atari with Deep Reinforcement Learning. Available at arXiv:1312.5602.
- [22] Mnih, V., Kavukcuoglu, K., Silver, D. et al. (2015). Human-Level Control through Deep Reinforcement Learning. *Nature*, 518, 529-533.
- [23] Mnih, V., Badia, A.P., Mirza, M., Graves, A., Lillicrap, T. P., Harley, T., Silver, D., & Kavukcuoglu, K. (2016). Asynchronous Method for Deep Reinforcement Learning. Available at arXiv:1602.01783.
- [24] Morse, P. M., & Kimball, G. E. (1951). *Methods of Operations Research*, New York: John Wiley and Sons.
- [25] Oroojlooyjadid, A., Nazari, M., Snyder, L., & Takac, M. (2020). A Deep Q-Network for the Beer Game: A Deep Reinforcement Learning Algorithm to Solve Inventory Optimization Problems. *Manufacturing & Service Operations Management*. Forthcoming. Available at arXiv: 1708.05924.
- [26] Peng, Z., Zhang, Y., Feng, Y., Zhang, T., Wu, Z., & Su, H. (2019). Deep Reinforcement Learning Approach for Capacitated Supply Chain Optimization under Demand Uncertainty. *IEEE Xplore 2019 Chinese Automation Congress*, 3512-3517.
- [27] Prak, D. & Teunter, R. (2019). A General Method for Addressing Forecasting Uncertainty in Inventory Models. *International Journal of Forecasting*, 35(1), 224-238.
- [28] Scarf, H. (1960). The Optimality of (S, s) Policies in the Dynamic Inventory Problem. In: Arrow, K., Karlin, S., & Suppes P. (eds), *Mathematical Methods in the Social Sciences*, Chapter 13. Stanford University Press, Stanford, CA.
- [29] Schulman, J., Levine, S., Moritz, P., Jordan, M. I., & Abbeel, P. (2015). Trust Region Policy Optimization. Available at arXiv: 1502.05477.
- [30] Schulman, J., Wolski, F., Dhariwal, P., Radford, A., & Klimov, O. (2017). Proximal Policy Optimization Algorithms. Available at arXiv: 1707.06347.
- [31] Silver, D., Lever, G., Heess, N., Degris, T., Wierstra, D., & Riedmiller, M. (2014). Deterministic Policy Gradient Algorithms. *Proceedings of the 31st International Conference on Machine Learning*, 32(1), 387-395.
- [32] Suetsugu, H., Narusue, Y., & Morikawa, H. (2019). Joint Replenishment Policy in Multi-Product Inventory System Using Branching Deep Q-Network with Reward Allocation. Available at Corpus: 201685454.
- [33] Vanvuchelen, N., Gijsbrechts, J., & Boute, R. N. (2020). Use of Proximal Policy Optimization for the Joint Replenishment Problem. *Computers in Industry*, 119, 103239.
- [34] Williams, R. J. (1992). Simple Statistical Gradient-Following Algorithms for Connectionist Reinforcement Learning. *Machine Learning*, 8, 229-256.



이 병 권

서울대학교 산업공학 학사
현재: 서울대학교 산업공학
석·박사 통합과정
관심 분야: 운영관리, 경영과학 응용



박 건 수

서울대학교 산업공학 학사
서울대학교 산업공학 석사
컬럼비아대학교 경영과학 박사
현재: 서울대학교 산업공학과 부교수
관심 분야: 공급사슬관리, 물류관리,
재고관리



정 세 윤

서울대학교 조정 · 지역시스템공학
& 연합전공 기술경영 학사
서울대학교 산업공학 석사
한국과학기술원 경영공학 박사
현재: 한국방송통신대학교 첨단공학부
산업공학전공 조교수
관심 분야: 공급사슬관리, 최적화 응용

한국SCM학회 연구 윤리 규정

제1조 (목적)

본 규정은 “한국SCM학회 연구 윤리 규정”이라 부르며 한국SCM학회(이하 “학회”라 한다)와 관련된 연구행위가 연구 목적을 달성하기 위해 수행되는 과정에서 인간의 기본적, 사회 공동 윤리를 손상하지 않도록 윤리 규정과 기준을 정함을 목적으로 한다. 여기서 연구 행위라는 것은 학회가 주관 공동 주관하는 학술대회와 학회 학술지와 관련된 연구 수행, 결과, 발표 및 게재 등을 포함한다.

제2조 (적용대상)

학회가 주관 또는 공동 주관하는 학술대회 발표와 학회 학술지 투고에 참여하는 학회의 회원들 외에 비회원들(이하 “저자”라 한다)에게도 준용된다.

제3조 (저자의 연구 윤리)

1. 저자는 아이디어의 도출, 실험 방법의 설계, 결과의 분석, 연구 결과의 발표, 연구 심사 등의 연구 행위에 있어 정직하여야 한다.
2. 저자는 타인의 연구나 주장의 전체 또는 일부분을 인용할 수 있다. 그러나 자신의 연구처럼 기술해서는 안 되며 반드시 정확하게 출처 표시와 참고문헌 목록을 작성하여야 한다.
3. 저자는 연구 수행과 결과에서 획득한 정보를 이용하여 부당한 이익을 추구하지 않는다.

제4조 (연구 내용의 기록 보존 및 공개)

1. 저자의 연구 내용은 타 연구자가 해석 및 확인이 용이하도록 정확하게 기록하여야 하며, 연구 수행 시 활용된 주요 사실 및 증거는 보존해야 한다.
2. 연구 결과가 출판된 후 타 연구자의 요청이 있을 경우 보안이 보장되는 범위 내에서 연구 결과물이 타 연구자의 연구 수행에 도움이 되도록 최대한 노력한다.

제5조 (저자의 책임과 보상)

1. 연구 결과에 기재된 모든 저자들은 발표된 사실에 책임을 다하도록 한다.
2. 저자는 공식적인 공동 연구자 또는 연구에 직간접적으로 기여한 사람들로만 구성되며 상대적 지위와 무관하게 학술적 기여도에 따라 저자 표기 순서가 결정된다.
3. 학회지 및 학술대회 발표논문집에 게재된 논문은 저자가 저작권을 가지나 공동의 목적으로 사용될 때는 학회가 사용권을 가진다.

제6조 (연구 부정행위)

연구 수행 중에 발생하는 부정행위는 다음과 같다.

1. 위조: 존재하지 않는 데이터나 연구결과를 허위로 만들어 내는 행위를 말한다.
2. 변조: 데이터의 변형이나 연구과정을 조작하여 연구결과를 왜곡하는 행위를 말한다.
3. 표절: 정당한 인용 없이 타 연구자의 연구 결과를 저자의 연구 결과에 사용하는 행위를 말한다.
4. 중복게재: 타 학술지에 게재 또는 투고 중인 원고를 본 학회지에 투고하는 행위를 말한다.
5. 부당한 논문 저자 표시: 연구 수행 중에 학술적 기여도가 없는 자에게 연구 결과의 저자 자격을 부여하는 행위를 말한다.

제7조 (윤리위원회 구성)

1. 학회는 연구 윤리와 관련된 사항을 검토·심의·의결하기 위해 학회 내에 윤리위원회를 운영한다.
2. 윤리위원회 구성은 위원장 1인과 부위원장 1인을 포함하여 5인으로 구성한다.
3. 윤리위원장은 학회 공동회장 중 한 분이 담당하며, 윤리위원회 부위원장은 학회지 공동 편집위원장 중 한 분을 윤리위원장이 임명하며, 나머지 3인의 위원회 회원은 윤리위원장과 부위원장의 합의로 임명한다.

제8조 (연구 부정행위 제재)

연구 부정행위가 적발된 연구 및 저자에 대해서는 윤리위원회의 검토를 거쳐 정도에 따라 다음과 같이 제재를 가할 수 있다.

1. 학회 징계 서한 발송
2. 학회의 해당 학회지에서 해당 연구 결과 삭제 또는 수정 요구
3. 연구 관련자의 적정 기간 동안 논문 투고 금지
4. 연구 관련자의 적정 기간 동안 회원자격 상실 및 연구 관련자 소속기관 세부사항 통보
5. 학회에서 제명

제9조 (윤리위원회 운영)

1. 필요한 연구 윤리 제정 및 개정을 담당한다.
2. 제소된 회원 및 연구에 대해 윤리 규정 위반 여부 심의 및 위반에 대한 제재를 의결한다.
3. 제소된 사안에 대해 접수된 날로부터 60일 이내에 심의·의결한다.
4. 위원회는 위원회의 조사 기간 동안 조사 내용 및 과정에 대해 일체의 보안을 유지하고, 관련자들의 신상정보를 보호한다.
5. 윤리위원회는 조사 결과 제소된 내용이 무혐의이거나 충분한 소명으로 혐의 사실이 해소될 경우 피고발자 혹은 혐의자의 명예를 회복하기 위해 적절한 후속 조치를 취할 수 있다.

제10조 (윤리위원회 제소 및 혐의자 의무)

1. 윤리위원회 제소는 회원 5인 이상의 서명을 받아야 한다.
2. 윤리위원회에 제소된 회원은 윤리위원회의 조사에 협조해야 한다.

제11조 (윤리위원회 의무)

1. 윤리위원회는 제소된 자에 대해 심의 결과가 확정되기 전까지는 회원으로 권리를 보장한다.
2. 윤리위원회에 제소된 자는 위원회에 충분히 소명할 권리를 갖으며, 위원회는 소명 및 반론 기회를 부여해야 한다.

제12조 기타 본 규정에 포함되지 않은 사항은 관계 법령과 사회적 규범에 의거 판단한다.**부 칙****제1조 (시행일)**

본 규정은 이사회에서 의결된 날부터 시행한다.

2013. 1. 16. 이사회 제정

Journal of the Korean Society of Supply Chain Management
Copyright Transfer Agreement

To: Editor of Journal of the Korean Society of Supply Chain Management

Title of submitted manuscript: _____

Author(s)(Full Names): _____

I hereby certify that I agreed to submit the manuscript entitled as above to Journal of the Korean Society of Supply Chain Management with the following statements:

- This manuscript is author's original work and has not been published before. It will not be submitted again to other journals without permission from Editor of Journal of the Korean Society of Supply Chain Management if it is accepted for publication.
- This manuscript should not contain any libelous statements, defamation and privacy intrusion. Any legal or ethical damage should not be directed to the Korean Society of Supply Chain Management due to this manuscript.
- All authors contributed to this manuscript have equal responsibility with respect to the copyright problem.
- Copyright of the manuscript to be published in the Journal of Korean Society of Supply Chain Management is transferred to the Korean Society of Supply Chain Management.

I agreed Declaration of Ethical Conduct in Research & Statement of Copyright Transfer.

Date:

Author(s) Name and Signature:

한국SCM학회지 제21권 제3호 심사자 명단(가나다 순)

김경민(명지대학교), 김기태(한밭대학교), 김보성(부산대학교), 김상국(한국과학기술정보연구원)
김용진(인하대학교), 김우제(서울과학기술대학교), 김창희(인천대학교), 민대기(이화여자대학교)
민윤홍(연세대학교), 서용원(중앙대학교), 이평수(경기대학교), 정병도(연세대학교)
정세윤(한국방송통신대학교), 정영선(전남대학교), 정태수(고려대학교), 최동현(한국항공대학교)

학회지 심사를 위해 노고를 아끼지 않은 심사자 여러분들께 깊은 감사의 말씀을 올립니다.

한국SCM학회지 제21권 제3호(특별호)

인 쇄 / 2021년 12월 31일

발 행 / 2021년 12월 31일

발행인 / 고창성

편집인 / 임성묵 · 박건수

발행처 / **한국SCM학회**

경기도 수원시 영통구 월드컵로 206 아주대학교
팔달관 812호

전화 031-211-5269

전송 031-214-5269

<http://www.kscm.org>

등록번호 ISSN 1598-382X