

Chapter 2

The Maximum Principle

This chapter represents the basic concepts of Classical Optimal Control related to the *Maximum Principle*. The formulation of the general optimal control problem in the Bolza (as well as in the Mayer and the Lagrange) form is presented. The Maximum Principle, which gives the necessary conditions of optimality, for various problems with a fixed and variable horizon is formulated and proven. All necessary mathematical claims are given in the Appendix, which makes this material self-contained.

This chapter is organized as follows. The classical optimal control problems in the Bolza, Lagrange, and Mayer form, are formulated in the next section. Then in Sect. 2.2 the variational inequality is derived based on the needle-shaped variations and Gronwall's inequality. Subsequently, a basic result is presented concerning the necessary conditions of the optimality for the problem considered in the Mayer form with terminal conditions using the duality relations.

2.1 Optimal Control Problem

2.1.1 Controlled Plant, Cost Functionals, and Terminal Set

Definition 2.1 Consider the controlled plant given by the following system of *Ordinary Differential Equations* (ODE):

$$\begin{cases} \dot{x}(t) = f(x(t), u(t), t), & \text{a.e. } t \in [0, T], \\ x(0) = x_0, \end{cases} \quad (2.1)$$

where $x = (x^1, \dots, x^n)^T \in \mathbb{R}^n$ is its state vector, and $u = (u^1, \dots, u^r)^T \in \mathbb{R}^r$ is the control that may run over a given control region $U \subset \mathbb{R}^r$ with the *cost functional*

$$J(u(\cdot)) := h_0(x(T)) + \int_{t=0}^T h(x(t), u(t), t) dt \quad (2.2)$$

containing the integral term as well as the terminal one, with the *terminal set* $\mathcal{M} \subseteq \mathbb{R}^n$ given by the inequalities

$$\mathcal{M} = \{x \in \mathbb{R}^n : g_l(x) \leq 0 \ (l = 1, \dots, L)\}. \quad (2.3)$$

The time process or *horizon* T is supposed to be fixed or variable and may be finite or infinite.

Definition 2.2 The function (2.2) is said to be given in the *Bolza form*. If in (2.2) $h_0(x) = 0$, we obtain the cost functional in the *Lagrange form*, that is,

$$J(u(\cdot)) = \int_{t=0}^T h(x(t), u(t), t) dt. \quad (2.4)$$

If in (2.2) $h(x, u, t) = 0$, we obtain the cost functional in the *Mayer form*, that is,

$$J(u(\cdot)) = h_0(x(T)). \quad (2.5)$$

Usually the following assumptions are assumed to be in force.

(A1) (U, d) is a separable metric space (with metric d) and $T > 0$.

(A2) The maps

$$\begin{cases} f : \mathbb{R}^n \times U \times [0, T] \rightarrow \mathbb{R}^n, \\ h : \mathbb{R}^n \times U \times [0, T] \rightarrow \mathbb{R}, \\ h_0 : \mathbb{R}^n \times U \times [0, T] \rightarrow \mathbb{R}, \\ g_l : \mathbb{R}^n \rightarrow \mathbb{R} \quad (l = 1, \dots, L) \end{cases} \quad (2.6)$$

are measurable and there exist a constant L and a continuity modulus $\bar{\omega} : [0, \infty) \rightarrow [0, \infty)$ such that for

$$\varphi = f(x, u, t), h(x, u, t), h_0(x, u, t), g_l(x) \ (l = 1, \dots, L)$$

the following inequalities hold:

$$\begin{cases} \|\varphi(x, u, t) - \varphi(\hat{x}, \hat{u}, t)\| \leq L\|x - \hat{x}\| + \bar{\omega}(d(u, \hat{u})) \\ \forall t \in [0, T], x, \hat{x} \in \mathbb{R}^n, u, \hat{u} \in U, \\ \|\varphi(0, u, t)\| \leq L \quad \forall u, t \in U \times [0, T]. \end{cases} \quad (2.7)$$

(A3) The maps

$$f, h, h_0 \text{ and } g_l \ (l = 1, \dots, L)$$

are of type C^1 in x and there exists a continuity modulus $\bar{\omega} : [0, \infty) \rightarrow [0, \infty)$ such that for

$$\varphi = f(x, u, t), h(x, u, t), h_0(x, u, t), g_l(x) \ (l = 1, \dots, L)$$

the following inequalities hold:

$$\left\| \frac{\partial}{\partial x} \varphi(x, u, t) - \frac{\partial}{\partial x} \varphi(\hat{x}, \hat{u}, t) \right\| \leq \bar{\omega}(\|x - \hat{x}\| + d(u, \hat{u}))$$

$$\forall t \in [0, T], x, \hat{x} \in \mathbb{R}^n, u, \hat{u} \in U. \quad (2.8)$$

2.1.2 Feasible and Admissible Control

Definition 2.3 A function $u(t)$, $t_0 \leq t \leq T$, is said to be a *feasible control* if it is measurable and $u(t) \in U$ for all $t \in [0, T]$. Denote the set of all feasible controls by

$$\mathcal{U}[0, T] := \{u(\cdot) : [0, T] \rightarrow U \mid u(t) \text{ is measurable}\}. \quad (2.9)$$

Definition 2.4 The control $u(t)$, $t_0 \leq t \leq T$ is said to be *admissible* if the terminal condition (2.3) holds, that is, if the corresponding trajectory $x(t)$ satisfies the terminal condition. We have the inclusion $x(T) \in \mathcal{M}$. Denote the set of all admissible controls by

$$\mathcal{U}_{\text{admis}}[0, T] := \{u(\cdot) : u(\cdot) \in \mathcal{U}[0, T], x(T) \in \mathcal{M}\}. \quad (2.10)$$

In view of the theorem on the existence of the solutions to the ODE (see Codrington and Levinson 1955 or Poznyak 2008), it follows that under the assumptions (A1)–(A2) for any $u(t) \in \mathcal{U}[0, T]$ (2.1) admits a unique solution, $x(\cdot) := x(\cdot, u(\cdot))$, and the functional (2.2) is well defined.

2.1.3 Setting of the Problem in the General Bolza Form

Based on the definitions given above, the optimal control problem (OCP) can be formulated as follows.

Problem 2.1 (OCP in the Bolza form)

$$\text{Minimize (2.2) over } \mathcal{U}_{\text{admis}}[0, T]. \quad (2.11)$$

Problem 2.2 (OCP with a fixed terminal term) If in problem (2.11)

$$\mathcal{M} = \{x_f \in \mathbb{R}^n\}$$

$$= \{x \in \mathbb{R}^n : g_1(x) = x - x_f \leq 0, g_2(x) = -(x - x_f) \leq 0\}$$

$$(\text{or, equivalently, } x = x_f) \quad (2.12)$$

then it is called an *optimal control problem with a fixed terminal term* x_f .

Definition 2.5 Any control $u^*(\cdot) \in \mathcal{U}_{\text{admis}}[0, T]$ satisfying

$$J(u^*(\cdot)) = \min_{u(\cdot) \in \mathcal{U}_{\text{admis}}[0, T]} J(u(\cdot)) \quad (2.13)$$

is called an *optimal control*, and the corresponding state trajectories $x^*(\cdot) := x^*(\cdot, u^*(\cdot))$ and $(x^*(\cdot), u^*(\cdot))$ are called an *optimal state trajectory* and an *optimal pair*.

2.1.4 Representation in the Mayer Form

Introduce the $(n + 1)$ -dimensional space $\mathbb{R}^{n+1}$ of the variables

$$x = (x_1, \dots, x_n, x_{n+1}),$$

where the first n coordinates satisfy (2.1) and the component x_{n+1} is given by

$$x_{n+1}(t) := \int_{\tau=0}^t h(x(\tau), u(\tau), \tau) d\tau \quad (2.14)$$

or, in differential form,

$$\dot{x}_{n+1}(t) = h(x(t), u(t), t) \quad (2.15)$$

with zero initial condition for the last component

$$x_{n+1}(0) = 0. \quad (2.16)$$

As a result, the initial Optimization Problem in the Bolza form (2.11) can be reformulated in the space $\mathbb{R}^{n+1}$ as a *Mayer Problem* with the cost functional $J(u(\cdot))$,

$$J(u(\cdot)) = h_0(x(T)) + x_{n+1}(T), \quad (2.17)$$

where the function $h_0(x)$ does not depend on the last coordinate $x_{n+1}(t)$, that is,

$$\frac{\partial}{\partial x_{n+1}} h_0(x) = 0. \quad (2.18)$$

From these relations it follows that the Mayer Problem with the cost function (2.17) is equivalent to the initial Optimization Control Problem (2.11) in the Bolza form.

2.2 Maximum Principle Formulation

2.2.1 Needle-Shaped Variations and Variational Equation

Let $(x^*(\cdot), u^*(\cdot))$ be the given optimal pair and $M_\varepsilon \subseteq [0, T]$ be a measurable set with the Lebesgue measure $|M_\varepsilon| = \varepsilon > 0$. Now let

$$u(\cdot) \in \mathcal{U}_{\text{admis}}[0, T]$$

be any given admissible control.

Definition 2.6 Define the control

$$u^\varepsilon(t) := \begin{cases} u^*(t) & \text{if } t \in [0, T] \setminus M_\varepsilon, \\ u(t) \in \mathcal{U}_{\text{admis}}[0, T] & \text{if } t \in M_\varepsilon. \end{cases} \quad (2.19)$$

It is evident that $u^\varepsilon(\cdot) \in \mathcal{U}_{\text{admis}}[0, T]$. In the following, $u^\varepsilon(\cdot)$ is referred to as a *needle-shaped* or *spike variation* of the optimal control $u^*(t)$.

The next lemma plays a key role in proving the basic MP-theorem. It gives an analytical estimation for the trajectories and for the cost function deviations. The corresponding differential equations can be interpreted also as “*sensitivity equations*.”

Lemma 2.1 (Variational equation) *Let $x^\varepsilon(\cdot) := x(\cdot, u^\varepsilon(\cdot))$ be the solution of (2.1) under the control $u^\varepsilon(\cdot)$ and let $\Delta^\varepsilon(\cdot)$ be the solution to the differential equation*

$$\begin{aligned} \dot{\Delta}^\varepsilon(t) &= \frac{\partial}{\partial x} f(x^*(t), u^*(t), t) \Delta^\varepsilon(t) \\ &\quad + [f(x^*(t), u^\varepsilon(t), t) - f(x^*(t), u^*(t), t)] \chi_{M_\varepsilon}(t), \\ \Delta^\varepsilon(0) &= 0, \end{aligned} \quad (2.20)$$

where $\chi_{M_\varepsilon}(t)$ is the characteristic function of the set M_ε , that is,

$$\chi_{M_\varepsilon}(t) := \begin{cases} 1 & \text{if } t \in M_\varepsilon, \\ 0 & \text{if } t \notin M_\varepsilon. \end{cases} \quad (2.21)$$

Then the following relations hold:

$$\begin{cases} \max_{t \in [0, T]} \|x^\varepsilon(t) - x^*(t)\| = O(\varepsilon), \\ \max_{t \in [0, T]} \|\Delta^\varepsilon(t)\| = O(\varepsilon), \\ \max_{t \in [0, T]} \|x^\varepsilon(t) - x^*(t) - \Delta^\varepsilon(t)\| = o(\varepsilon), \end{cases} \quad (2.22)$$

and the following variational equations hold:

(a) for the cost function given in the Bolza form (2.2)

$$\begin{aligned} &J(u^\varepsilon(\cdot)) - J(u^*(\cdot)) \\ &= \left(\frac{\partial}{\partial x} h_0(x^*(T)), \Delta^\varepsilon(T) \right) \\ &\quad + \int_{t=0}^T \left\{ \left(\frac{\partial}{\partial x} h(x^*(t), u^*(t), t), \Delta^\varepsilon(t) \right) \right. \\ &\quad \left. + [h(x^*(t), u^\varepsilon(t), t) - h(x^*(t), u^*(t), t)] \chi_{M_\varepsilon}(t) \right\} dt \\ &\quad + o(\varepsilon) \end{aligned} \quad (2.23)$$

(b) for the cost function given in the Mayer form (2.5)

$$J(u^\varepsilon(\cdot)) - J(u^*(\cdot)) = \left(\frac{\partial}{\partial x} h_0(x^*(T)), \Delta^\varepsilon(T) \right) + o(\varepsilon) \quad (2.24)$$

Proof Define

$$\delta^\varepsilon(t) := x^\varepsilon(t) - x^*(t).$$

Then by assumption (A2) (2.7) for any $t \in [0, T]$ it follows that

$$\|\delta^\varepsilon(t)\| \leq \int_{s=0}^t L \|\delta^\varepsilon(s)\| ds + K\varepsilon, \quad (2.25)$$

which, by Gronwall's Lemma (see Appendix 2.3 of this chapter), implies the first relation in (2.22). Define

$$\eta^\varepsilon(t) := x^\varepsilon(t) - x^*(t) - \Delta^\varepsilon(t) = \delta^\varepsilon(t) - \Delta^\varepsilon(t). \quad (2.26)$$

Then we have

$$\begin{aligned} \dot{\eta}^\varepsilon(t) &= [f(x^\varepsilon(t), u^\varepsilon(t), t) - f(x^*(t), u^*(t), t)] \\ &\quad - \frac{\partial}{\partial x} f(x^*(t), u^*(t), t) \Delta^\varepsilon(t) \\ &\quad - [f(x^*(t), u^\varepsilon(t), t) - f(x^*(t), u^*(t), t)] \chi_{M_\varepsilon}(t) \\ &= \int_{\theta=0}^1 \frac{\partial}{\partial x} f(x^*(t) + \theta \delta^\varepsilon(t), u^\varepsilon(t), t) d\theta \cdot \delta^\varepsilon(t) \\ &\quad - \frac{\partial}{\partial x} f(x^*(t), u^*(t), t) [\delta^\varepsilon(t) - \eta^\varepsilon(t)] \\ &\quad - [f(x^*(t), u^\varepsilon(t), t) - f(x^*(t), u^*(t), t)] \chi_{M_\varepsilon}(t) \\ &= \int_{\theta=0}^1 \left[\frac{\partial}{\partial x} f(x^*(t) + \theta \delta^\varepsilon(t), u^\varepsilon(t), t) \right. \\ &\quad \left. - \frac{\partial}{\partial x} f(x^*(t), u^*(t), t) d\theta \right] \delta^\varepsilon(t) \\ &\quad - [f(x^*(t), u^\varepsilon(t), t) - f(x^*(t), u^*(t), t)] \chi_{M_\varepsilon}(t) \\ &\quad + \frac{\partial}{\partial x} f(x^*(t), u^*(t), t) \eta^\varepsilon(t). \end{aligned} \quad (2.27)$$

Integrating the last identity (2.27) and holding in view (A2) (2.7) and (A3) (2.8), we obtain

$$\begin{aligned} \|\eta^\varepsilon(t)\| &\leq \int_{s=0}^t \int_{\theta=0}^1 [\bar{\omega}(\theta \|\delta^\varepsilon(s)\| + d(u^\varepsilon(s), u^*(s)))] \|\delta^\varepsilon(s)\| d\theta ds \\ &\quad + \int_{s=0}^t \bar{\omega}(d(u^\varepsilon(s), u^*(s))) \chi_{M_\varepsilon}(s) ds \end{aligned}$$

$$\begin{aligned}
& + \int_{s=0}^t \frac{\partial}{\partial x} f(x^*(s), u^*(s), s) \eta^\varepsilon(s) \, ds \\
& \leq \varepsilon o(1) + \text{Const} \int_{s=0}^t \|\eta^\varepsilon(s)\| \, ds.
\end{aligned} \tag{2.28}$$

The last inequality in (2.28) by Gronwall's Lemma directly implies the third relation in (2.22). The second relation is a consequence of the first and third ones. The same manipulations lead to (2.23) and (2.24). $\square$

2.2.2 Adjoint Variables and MP Formulation for Cost Functionals with a Fixed Horizon

The classical format of the MP formulation gives a set of *first-order necessary conditions* for the optimal pairs.

Theorem 2.1 (MP for the Mayer Form with a fixed horizon) *If under the assumptions (A1)–(A3) a pair $(x^*(\cdot), u^*(\cdot))$ is optimal, then there exist vector functions $\psi(t)$ satisfying the system of adjoint equations*

$$\dot{\psi}(t) = -\frac{\partial}{\partial x} f(x^*(t), u^*(t), t)^T \psi(t) \quad \text{a.e. } t \in [0, T] \tag{2.29}$$

and nonnegative constants $\mu \geq 0$ and $v_l \geq 0$ ($l = 1, \dots, L$) such that the following four conditions hold.

1. (The maximality condition) *For almost all $t \in [0, T]$*

$$H(\psi(t), x^*(t), u^*(t), t) = \max_{u \in U} H(\psi(t), x^*(t), u, t), \tag{2.30}$$

where the Hamiltonian is defined as

$$H(\psi, x, u, t) := \psi^T f(x, u, t), \quad t, x, u, \psi \in [0, T] \times \mathbb{R}^n \times \mathbb{R}^r \times \mathbb{R}^n. \tag{2.31}$$

2. (Transversality condition) *The equality*

$$\psi(T) + \mu \frac{\partial}{\partial x} h_0(x^*(T)) + \sum_{l=1}^L v_l \frac{\partial}{\partial x} g_l(x^*(T)) = 0 \tag{2.32}$$

holds.

3. (Complementary slackness conditions) *Either the equality*

$$g_l(x^*(T)) = 0$$

holds, or $v_l = 0$, that is, for any $l = 1, \dots, L$,

$$v_l g_l(x^*(T)) = 0. \tag{2.33}$$

4. (Nontriviality condition) *At least one of the quantities $|\psi(T)|$ and v_l is distinct from zero, that is,*

$$|\psi(T)| + \mu + \sum_{l=1}^L v_l > 0. \quad (2.34)$$

Proof Let $\psi(t)$ be the solution to (2.29) corresponding to the terminal condition $\psi(T) = b$ and $\bar{t} \in [0, T]$. Define $M_\varepsilon := [\bar{t}, \bar{t} + \varepsilon] \subseteq [0, T]$. If $u^*(t)$ is an optimal control, then according to the Lagrange Principle¹ there exist constants $\mu \geq 0$ and $v_l \geq 0$ ($l = 1, \dots, L$) such that for any $\varepsilon \geq 0$

$$\mathcal{L}(u^\varepsilon(\cdot), \mu, v) - \mathcal{L}(u^*(\cdot), \mu, v) \geq 0. \quad (2.35)$$

Here

$$\mathcal{L}(u(\cdot), \mu, v) := \mu J(u(\cdot)) + \sum_{l=1}^L v_l g_l(x(T)). \quad (2.36)$$

Taking into account that

$$\psi(T) = b$$

and

$$\Delta^\varepsilon(0) = 0$$

by the differential chain rule, applied to the term $\psi(t)^T \Delta^\varepsilon(t)$, and in view of (2.20) and (2.29), we obtain

$$\begin{aligned} b^T \Delta^\varepsilon(T) &= \psi(T)^T \Delta^\varepsilon(T) - \psi(0)^T \Delta^\varepsilon(0) \\ &= \int_{t=0}^T d(\psi(t)^T \Delta^\varepsilon(t)) \\ &= \int_{t=0}^T (\dot{\psi}(t)^T \Delta^\varepsilon(t) + \psi(t)^T \dot{\Delta}^\varepsilon(t)) dt \\ &= \int_{t=0}^T \left[-\Delta^\varepsilon(t)^T \frac{\partial}{\partial x} f(x^*(t), u^*(t), t)^T \psi(t) \right. \\ &\quad \left. + \psi(t)^T \frac{\partial}{\partial x} f(x^*(t), u^*(t), t) \Delta^\varepsilon(t) \right. \\ &\quad \left. + \psi(t)^T [f(x^*(t), u^\varepsilon(t), t) - f(x^*(t), u^*(t), t)] \chi_{M_\varepsilon}(t) \right] dt \\ &= \int_{t=0}^T \psi(t)^T [f(x^*(t), u^\varepsilon(t), t) - f(x^*(t), u^*(t), t)] \chi_{M_\varepsilon}(t) dt. \end{aligned} \quad (2.37)$$

¹See Appendix 2.3 for finite-dimensional spaces.

The variational equality (2.23) together with (2.35) and (2.37) implies

$$\begin{aligned}
0 &\leq \mathcal{L}(u^\varepsilon(\cdot), \mu, v) - \mathcal{L}(u^*(\cdot), \mu, v) \\
&= \mu \left(\frac{\partial}{\partial x} h_0(x^*(T)), \Delta^\varepsilon(T) \right) + b^T \Delta^\varepsilon(T) \\
&\quad - \int_{t=0}^T \psi(t)^T (f(x^*(t), u^\varepsilon(t), t) - f(x^*(t), u^*(t), t)) \chi_{M_\varepsilon}(t) dt \\
&\quad + \sum_{l=1}^L v_l [g_l(x(T)) - g_l(x^*(T))] + o(\varepsilon) \\
&= \left(\mu \frac{\partial}{\partial x} h_0(x^*(T)) + b + \sum_{l=1}^L v_l \frac{\partial}{\partial x} g_l(x^*(T)), \Delta^\varepsilon(T) \right) \\
&\quad - \int_{t=\bar{t}}^{\bar{t}+\varepsilon} [\psi(t)^T (f(x^*(t), u^\varepsilon(t), t) - f(x^*(t), u^*(t), t))] dt + o(\varepsilon) \\
&= \left(\mu \frac{\partial}{\partial x} h_0(x^*(T)) + b + \sum_{l=1}^L v_l \frac{\partial}{\partial x} g_l(x^*(T)), \Delta^\varepsilon(T) \right) \\
&\quad - \int_{t=\bar{t}}^{\bar{t}+\varepsilon} [H(\psi(t), x^*(t), u^\varepsilon(t), t) - H(\psi(t), x^*(t), u^*(t), t)] dt. \quad (2.38)
\end{aligned}$$

(1) Letting ε go to zero from (2.38) it follows that

$$0 \leq \left(\mu \frac{\partial}{\partial x} h_0(x^*(T)) + b + \sum_{l=1}^L v_l \frac{\partial}{\partial x} g_l(x^*(T)), \Delta^\varepsilon(T) \right) \Big|_{\varepsilon=0},$$

which should be valid for any $\Delta^\varepsilon(T)|_{\varepsilon=0}$. This is possible only if (this can be proved by the construction)

$$\mu \frac{\partial}{\partial x} h_0(x^*(T)) + b + \sum_{l=1}^L v_l \frac{\partial}{\partial x} g_l(x^*(T)) = 0, \quad (2.39)$$

which is equivalent to (2.32). Thus, the transversality condition is proven.

(2) In view of (2.39) the inequality (2.38) is simplified to

$$0 \leq - \int_{t=\bar{t}}^{\bar{t}+\varepsilon} [H(\psi(t), x^*(t), u^\varepsilon(t), t) - H(\psi(t), x^*(t), u^*(t), t)] dt. \quad (2.40)$$

This inequality together with the separability of the metric space U directly leads to the Maximality Condition (2.30).

(3) Suppose that (2.33) does not hold, that is, there exist an index l_0 and a multiplier $\tilde{v}_{l_0}$ such that

$$v_l g_l(x^*(T)) < 0.$$

This implies that

$$\begin{aligned} \mathcal{L}(u^*(\cdot), \mu, \tilde{v}) &:= \mu J(u^*(\cdot)) + \sum_{l=1}^L \tilde{v}_l g_l(x^*(T)) \\ &= \mu J(u^*(\cdot)) + \tilde{v}_{l_0} g_{l_0}(x^*(T)) \\ &< \mu J(u^*(\cdot)) = \mathcal{L}(u^*(\cdot), \mu, v). \end{aligned}$$

It means that $u^*(\cdot)$ is not an optimal control. We obtain a contradiction. So the complementary slackness condition is proven.

(4) Suppose that (2.34) is not valid, that is,

$$|\psi(T)| + \mu + \sum_{l=1}^L v_l = 0,$$

which implies

$$\psi(T) = 0, \quad \mu = v_l = 0 \quad (l = 1, \dots, L)$$

and, hence, in view of (2.29) and Gronwall's Lemma it follows that $\psi(t) = 0$ for all $t \in [0, T]$. So,

$$H(\psi(t), x(t), u(t), t) = 0$$

for any $u(t)$ (not only for $u^*(t)$). This means that the application of any admissible control keeps the cost function unchanged and this corresponds to the trivial situation of an “uncontrollable” system. So the nontriviality condition is proven as well. $\square$

2.2.3 The Regular Case

In the so-called *regular case*, when $\mu > 0$ (this means that the nontriviality condition holds automatically), the variable $\psi(t)$ and constants v_l may be normalized and changed to $\tilde{\psi}(t) := \psi(t)/\mu$ and $\tilde{v}_l := v_l/\mu$. In this new variable the MP formulation looks as follows.

Theorem 2.2 (MP in the regular case) *If under the assumptions (A1)–(A3) a pair $(x^*(\cdot), u^*(\cdot))$ is optimal then there exist vector functions $\tilde{\psi}(t)$ satisfying the system of adjoint equations*

$$\frac{d}{dt} \tilde{\psi}(t) = -\frac{\partial}{\partial x} f(x^*(t), u^*(t), t)^T \tilde{\psi}(t) \quad \text{a.e. } t \in [0, T]$$

and $v_l \geq 0$ ($l = 1, \dots, L$) such that the following three conditions hold.

1. (The maximality condition) For almost all $t \in [0, T]$

$$H(\tilde{\psi}(t), x^*(t), u^*(t), t) = \max_{u \in U} H(\tilde{\psi}(t), x^*(t), u, t),$$

where the Hamiltonian is defined as

$$H(\psi, x, u, t) := \tilde{\psi}^T f(x, u, t), \quad t, x, u, \psi \in [0, T] \times \mathbb{R}^n \times \mathbb{R}^r \times \mathbb{R}^n.$$

2. (Transversality condition) For every $\alpha \in \mathcal{A}$, the equalities

$$\tilde{\psi}(T) + \frac{\partial}{\partial x} h_0(x^*(T)) + \sum_{l=1}^L \tilde{v}_l \frac{\partial}{\partial x} g_l(x^*(T)) = 0$$

hold.

3. (Complementary slackness conditions) Either the equality

$$g_l(x^*(T)) = 0$$

holds, or

$$v_l = 0,$$

that is, for any $l = 1, \dots, L$

$$v_l g_l(x^*(T)) = 0.$$

Remark 2.1 This means that without loss of generality we may put $\mu = 1$. It may be shown (Polyak 1987) that the regularity property holds if the vectors $\frac{\partial}{\partial x} g_l(x^*(T))$ are linearly independent. The verification of this property is usually not so simple a task. There are also other known weaker conditions of regularity (Poznyak 2008) (Slater's condition, etc.).

2.2.4 Hamiltonian Form and Constancy Property

Corollary 2.1 (Hamiltonian for the Bolza Problem) *The Hamiltonian (2.31) for the Bolza Problem has the form*

$$\begin{aligned} H(\psi, x, u, t) &:= \psi^T f(x, u, t) - \mu h(x(t), u(t), t), \\ t, x, u, \psi &\in [0, T] \times \mathbb{R}^n \times \mathbb{R}^r \times \mathbb{R}^n. \end{aligned} \quad (2.41)$$

Proof This follows from (2.14)–(2.18). Indeed, the representation in the Mayer form,

$$\dot{x}_{n+1}(t) = h(x(t), u(t), t),$$

implies

$$\dot{\psi}_{n+1}(t) = 0$$

and, hence,

$$\psi_{n+1}(T) = -\mu. \quad \square$$

Corollary 2.2 (Hamiltonian form) *Equations (2.1) and (2.29) may be represented in the so-called Hamiltonian form (forward–backward ODE form):*

$$\left\{ \begin{array}{l} \dot{x}^*(t) = \frac{\partial}{\partial \psi} H(\psi(t), x^*(t), u^*(t), t), \quad x^*(0) = x_0, \\ \dot{\psi}(t) = -\frac{\partial}{\partial x} H(\psi(t), x^*(t), u^*(t), t), \\ \psi(T) = -\mu \frac{\partial}{\partial x} h_0(x^*(T)) - \sum_{l=1}^L v_l \frac{\partial}{\partial x} g_l(x^*(T)). \end{array} \right. \quad (2.42)$$

Proof It directly follows from the comparison of the right-hand side of (2.31) with (2.1) and (2.29). $\square$

Corollary 2.3 (Constancy property) *For stationary systems, see (2.1)–(2.2), where*

$$f = f(x(t), u(t)), \quad h = h(x(t), u(t)). \quad (2.43)$$

It follows that for all $t \in [t_0, T]$

$$H(\psi(t), x^*(t), u^*(\psi(t), x^*(t))) = \text{const}. \quad (2.44)$$

Proof One can see that in this case

$$H = H(\psi(t), x(t), u(t)),$$

that is,

$$\frac{\partial}{\partial t} H = 0.$$

Hence, $u^*(t)$ is a function of $\psi(t)$ and $x^*(t)$ only, that is,

$$u^*(t) = u^*(\psi(t), x^*(t)).$$

Denote

$$H(\psi(t), x^*(t), u^*(\psi(t), x^*(t))) := \tilde{H}(\psi(t), x^*(t)).$$

Then (2.42) becomes

$$\left\{ \begin{array}{l} \dot{x}(t) = \frac{\partial}{\partial \psi} \tilde{H}(\psi(t), x^*(t)), \\ \dot{\psi}(t) = -\frac{\partial}{\partial x} \tilde{H}(\psi(t), x^*(t)), \end{array} \right.$$

which implies

$$\begin{aligned} \frac{d}{dt} \tilde{H}(\psi(t), x^*(t)) &= \frac{\partial}{\partial \psi} \tilde{H}(\psi(t), x^*(t))^T \dot{\psi}(t) \\ &+ \frac{\partial}{\partial x} \tilde{H}(\psi(t), x^*(t))^T \dot{x}(t) = 0 \end{aligned}$$

and hence for any $t \in [t_0, T]$

$$\tilde{H}(\psi(t), x^*(t)) = \text{const.} \quad (2.45)$$

□

2.2.5 Variable Horizon Optimal Control Problem and Zero Property

Consider now the following generalization of the optimal control problem (2.1), (2.5), (2.11) permitting the terminal time to be free. In view of this, the optimization problem may be formulated in the following manner:

$$\begin{aligned} &\text{minimize } J(u(\cdot)) = h_0(x(T), T) \\ &\text{over } u(\cdot) \in \mathcal{U}_{\text{admis}}[0, T] \end{aligned} \quad (2.46)$$

and $T \geq 0$ with the terminal set $\mathcal{M}(T)$ given by

$$\mathcal{M}(T) = \{x(T) \in \mathbb{R}^n : g_l(x(T), T) \leq 0 \ (l = \overline{1, L})\}. \quad (2.47)$$

Theorem 2.3 (MP for the variable horizon case) *If under the assumptions (A1)–(A3) the pair $(T^*, u^*(\cdot))$ is a solution of the problem (2.46)–(2.47) and $x^*(t)$ is the corresponding optimal trajectory, then there exist vector functions $\psi(t)$ satisfying the system of adjoint equations (2.29) and nonnegative constants*

$$\mu \geq 0 \quad \text{and} \quad v_l \geq 0 \quad (l = \overline{1, L})$$

such that all four conditions of Theorem 2.1 are fulfilled and, in addition, the following condition for the terminal time holds:

$$\begin{aligned} H(\psi(T), x(T), u(T), T) &:= \psi^T(T) f(x(T), u(T - 0), T) \\ &= \mu \frac{\partial}{\partial T} h_0(x^*(T), T) + \sum_{l=1}^L v_l \frac{\partial}{\partial T} g_l(x^*(T), T). \end{aligned} \quad (2.48)$$

Proof Since $(T^*, u^*(\cdot))$ is a solution of the problem, evidently $u^*(\cdot)$ is a solution of the problem (2.1), (2.5), (2.11) with the fixed horizon $T = T^*$ and, hence, all four properties of Theorem 2.1 with $T = T^*$ should be fulfilled. Let us find the additional condition to the terminal time T^* which should also be satisfied.

(a) Consider again, as in (2.19), the needle-shaped variation defined by

$$u^\varepsilon(t) := \begin{cases} u^*(t) & \text{if } t \in [0, T^*] \setminus (M_\varepsilon \wedge (T^* - \varepsilon, T^*]), \\ u(t) \in \mathcal{U}_{\text{admis}}[0, T^*] & \text{if } t \in M_\varepsilon \subseteq [0, T^* - \varepsilon), \\ u(t) \in \mathcal{U}_{\text{admis}}[0, T^*] & \text{if } t \in [T^* - \varepsilon, T^*]. \end{cases} \quad (2.49)$$

Then, for $\mathcal{L}(u(\cdot), \mu, v, T)$ defined as

$$\mathcal{L}(u(\cdot), \mu, v, T) := \mu J(u(\cdot), T) + \sum_{l=1}^L v_l g_l(x(T), T), \quad (2.50)$$

it follows that

$$\begin{aligned} 0 &\leq \mathcal{L}(u^\varepsilon(\cdot), \mu, v, T^* - \varepsilon) - \mathcal{L}(u^*(\cdot), \mu, v, T^*) \\ &= \mu h_0(x(T^* - \varepsilon), T^* - \varepsilon) + \sum_{l=1}^L v_l g_l(x(T^* - \varepsilon), T^* - \varepsilon) \\ &\quad - \mu h_0(x^*(T^*), T^*) - \sum_{l=1}^L v_l g_l(x^*(T^*), T^*). \end{aligned}$$

Hence, by applying the transversality condition (2.32) we obtain

$$\begin{aligned} 0 &\leq -\varepsilon \left(\mu \frac{\partial}{\partial T} h_0(x(T^*), T^*) + v_l \frac{\partial}{\partial T} g_l(x(T^*), T^*) \right) \\ &\quad + o(\varepsilon) - \varepsilon \left(\mu \frac{\partial}{\partial x} h_0(x(T^*), T^*) \right. \\ &\quad \left. + \sum_{l=1}^L v_l \frac{\partial}{\partial x} g_l(x(T^*), T^*), f(x(T^*), u^*(T^* - 0), T^*) \right) \\ &= -\varepsilon \left(\mu \frac{\partial}{\partial T} h_0(x(T^*), T^*) + v_l \frac{\partial}{\partial T} g_l(x(T^*), T^*) \right) \\ &\quad + \varepsilon \psi^T(T^*) f(x(T^*), u^*(T^* - 0), T^*) + o(\varepsilon) \\ &= -\varepsilon \left(\mu \frac{\partial}{\partial T} h_0(x(T^*), T^*) + v_l \frac{\partial}{\partial T} g_l(x(T^*), T^*) \right) \\ &\quad + \varepsilon H(\psi(T^*), x^*(T^*), u^*(T - 0), T^*) + o(\varepsilon), \end{aligned}$$

which, by dividing by ε and letting ε go to zero, implies

$$\begin{aligned} &H(\psi(T^*), x^*(T^*), u^*(T - 0), T^*) \\ &\geq \mu \frac{\partial}{\partial T} h_0(x(T^*), T^*) + v_l \frac{\partial}{\partial T} g_l(x(T^*), T^*). \end{aligned} \quad (2.51)$$

(b) Analogously, for the needle-shaped variation

$$u^\varepsilon(t) := \begin{cases} u^*(t) & \text{if } t \in [0, T^*] \setminus M_\varepsilon, \\ u(t) \in \mathcal{U}_{\text{admis}}[0, T^*] & \text{if } t \in M_\varepsilon, \\ u^*(T^* - 0) & \text{if } t \in [T^*, T^* + \varepsilon], \end{cases} \quad (2.52)$$

it follows that

$$\begin{aligned} 0 &\leq \mathcal{L}(u^\varepsilon(\cdot), \mu, v, T^* + \varepsilon) - \mathcal{L}(u^*(\cdot), \mu, v, T^*) \\ &= \varepsilon \left(\mu \frac{\partial}{\partial T} h_0(x(T^*), T^*) + v_l \frac{\partial}{\partial T} g_l(x(T^*), T^*) \right) \\ &\quad - \varepsilon H(\psi(T^*), x^*(T^*), u^*(T - 0), T^*) + o(\varepsilon) \end{aligned}$$

and

$$\begin{aligned} &H(\psi(T^*), x^*(T^*), u^*(T - 0), T^*) \\ &\leq \mu \frac{\partial}{\partial T} h_0(x(T^*), T^*) + v_l \frac{\partial}{\partial T} g_l(x(T^*), T^*). \end{aligned} \quad (2.53)$$

Combining (2.49) and (2.52), we obtain (2.48). The theorem is proven. $\square$

Corollary 2.4 (Zero property) *If under the conditions of the theorem above the functions*

$$h_0(x, T), \quad g_l(x, T) \quad (l = 1, \dots, L)$$

do not depend on T directly, that is,

$$\frac{\partial}{\partial T} h_0(x, T) = \frac{\partial}{\partial T} g_l(x, T) = 0 \quad (l = 1, \dots, L)$$

then

$$H(\psi(T^*), x^*(T^*), u^*(T - 0), T^*) = 0. \quad (2.54)$$

If, in addition, the stationary case is considered (see (2.43)), then (2.54) holds for all $t \in [0, T^]$, that is,*

$$H(\psi(t), x^*(t), u^*(\psi(t), x^*(t))) = 0. \quad (2.55)$$

Proof The result directly follows from (2.44) and (2.54). $\square$

2.2.6 Joint Optimal Control and Parametric Optimization Problem

Consider the nonlinear plant given by

$$\begin{cases} \dot{x}_a(t) = f(x_a(t), u(t), t; a), & \text{a.e. } t \in [0, T], \\ x_a(0) = x_0 \end{cases} \quad (2.56)$$

at the fixed horizon T , where $a \in \mathbb{R}^p$ is a vector of parameters that also can be selected to optimize the functional (2.5), which in this case is

$$J(u(\cdot), a) = h_0(x_a(T)). \quad (2.57)$$

(A4) It will be supposed that the right-hand side of (2.56) is differentiable for all $a \in \mathbb{R}^p$.

In view of this, OCP is formulated as

$$\begin{aligned} & \text{minimize } J(u(\cdot), a) \quad (2.57) \\ & \text{over } \mathcal{U}_{\text{admis}}[0, T] \text{ and } a \in \mathbb{R}^p. \end{aligned} \quad (2.58)$$

Theorem 2.4 (Joint OC and parametric optimization) *If under the assumptions (A1)–(A4) the pair $(u^*(\cdot), a^*)$ is a solution of the problem (2.46)–(2.47) and $x^*(t)$ is the corresponding optimal trajectory, then there exist vector functions $\psi(t)$ satisfying the system of the adjoint equations (2.29) with*

$$x^*(t), u^*(t), a^*$$

and nonnegative constants

$$\mu \geq 0 \quad \text{and} \quad v_l \geq 0 \quad (l = 1, \dots, L)$$

such that all four conditions of Theorem 2.1 are fulfilled and, in addition, the following condition for the optimal parameter holds:

$$\int_{t=0}^T \frac{\partial}{\partial a} H(\psi(t), x^*(t), u^*(t), t; a^*) dt = 0. \quad (2.59)$$

Proof For this problem $\mathcal{L}(u(\cdot), \mu, v, a)$ is defined as previously:

$$\mathcal{L}(u(\cdot), \mu, v, a) := \mu h_0(x(T)) + \sum_{l=1}^L v_l g_l(x(T)). \quad (2.60)$$

Introduce the matrix

$$\Delta^a(t) = \frac{\partial}{\partial a} x^*(t) \in \mathbb{R}^{n \times p},$$

called the *matrix of sensitivity (with respect to parameter variations)*, which satisfies the following differential equation:

$$\begin{aligned} \dot{\Delta}^a(t) &= \frac{d}{dt} \frac{\partial}{\partial a} x^*(t) = \frac{\partial}{\partial a} \dot{x}^*(t) \\ &= \frac{\partial}{\partial a} f(x^*(t), u^*(t), t; a^*) \\ &= \frac{\partial}{\partial a} f(x^*(t), u^*(t), t; a^*) + \frac{\partial}{\partial x} f(x^*(t), u^*(t), t; a^*) \Delta^a(t), \\ \Delta^a(0) &= 0. \end{aligned} \quad (2.61)$$

In view of this and using (2.29), it follows that

$$\begin{aligned}
0 &\leq \mathcal{L}(u^*(\cdot), \mu, \nu, a) - \mathcal{L}(u^*(\cdot), \mu, \nu, a^*) \\
&= (a - a^*)^T \Delta^a(T)^T \left(\mu \frac{\partial}{\partial x} h_0(x^*(T)) + \sum_{l=1}^L \nu_l \frac{\partial}{\partial x} g_l(x^*(T)) \right) + o(\|a - a^*\|) \\
&= (a - a^*)^T \Delta^a(T)^T \psi(T) + o(\|a - a^*\|) \\
&= (a - a^*)^T [\Delta^a(T)^T \psi(T) - \Delta^a(0)^T \psi(0)] + o(\|a - a^*\|) \\
&= (a - a^*)^T \int_{t=0}^T d[\Delta^a(t)^T \psi(t)] + o(\|a - a^*\|) \\
&= (a - a^*)^T \int_{t=0}^T \left[-\Delta^a(t)^T \frac{\partial}{\partial x} f(x^*(t), u^*(t), t; a^*)^T \psi(t) \right. \\
&\quad \left. + \Delta^a(t)^T \frac{\partial}{\partial x} f(x^*(t), u^*(t), t; a^*)^T \psi(t) \right. \\
&\quad \left. + \frac{\partial}{\partial a} f(x^*(t), u^*(t), t; a^*)^T \psi(t) \right] dt + o(\|a - a^*\|) \\
&= (a - a^*)^T \int_{t=0}^T \frac{\partial}{\partial a} f(x^*(t), u^*(t), t; a^*)^T \psi(t) dt + o(\|a - a^*\|).
\end{aligned}$$

But this inequality is possible for any $a \in \mathbb{R}^p$ in a small neighborhood of a^* only if (2.59) holds (this may be proved by contradiction). The theorem is proven. $\square$

2.2.7 Sufficient Conditions of Optimality

Some additional notions and constructions related to *Convex Analysis* will be useful later on. Let $\partial F(x)$ be a *subgradient* convex (not necessarily differentiable) function $F(x)$ at $x \in \mathbb{R}^n$, that is, $\forall y \in \mathbb{R}^n$

$$\partial F(x) := \{a \in \mathbb{R}^n : F(x+y) \geq F(x) + (a, y)\}. \quad (2.62)$$

Lemma 2.2 (Criterion of Optimality) *The condition*

$$0 \in \partial F(x^*) \quad (2.63)$$

is necessary and sufficient for guaranteeing that x^ is a solution to the finite-dimensional optimization problem*

$$\min_{x \in X \subseteq \mathbb{R}^n} F(x). \quad (2.64)$$

Proof (a) *Necessity.* Let x^* be one of the points minimizing $F(x)$ in $X \subseteq \mathbb{R}^n$. Then for any $y \in X$

$$F(x^* + y) \geq F(x^*) + (0, y) = F(x^*),$$

which means that 0 is a subgradient $F(x)$ at the point x^* .

(b) *Sufficiency.* If 0 is a subgradient $F(x)$ at the point x^* , then

$$F(x^* + y) \geq F(x^*) + (0, y) = F(x^*)$$

for any $y \in X$, which means that the point x^* is a solution of the problem (2.64). $\square$

An additional assumption concerning the control region is also required.

(A5) The control domain U is supposed to be a *convex body* (that is, it is convex and has a nonempty interior).

Lemma 2.3 (On mixed subgradient) *Let φ be a convex (or concave) function on $\mathbb{R}^n \times U$ where U is a convex body. Assuming that $\varphi(x, u)$ is differential in x and is continuous in (x, u) , the following inclusion turns out to be valid for any $(x^*, u^*) \in \mathbb{R}^n \times U$:*

$$\{(\varphi_x(x^*, u^*), r) : r \in \partial_u \varphi(x^*, u^*)\} \subseteq \partial_{x,u} \varphi(x^*, u^*). \quad (2.65)$$

Proof For any $y \in \mathbb{R}^n$ in view of the convexity of φ and its differentiability on x , it follows that

$$\varphi(x^* + y, u^*) - \varphi(x^*, u^*) \geq (\varphi_x(x^*, u^*), y). \quad (2.66)$$

Similarly, in view of the convexity of φ in u , there exists a vector $r \in \mathbb{R}^r$ such that for any $x^*, y \in \mathbb{R}^n$ and any $\bar{u} \in U$

$$\varphi(x^* + y, u^* + \bar{u}) - \varphi(x^* + y, u^*) \geq (r, \bar{u}). \quad (2.67)$$

So, taking into account the previous inequalities (2.66)–(2.67), we obtain

$$\begin{aligned} \varphi(x^* + y, u^* + \bar{u}) - \varphi(x^*, u^*) &= [\varphi(x^* + y, u^* + \bar{u}) - \varphi(x^* + y, u^*)] \\ &\quad + [\varphi(x^* + y, u^*) - \varphi(x^*, u^*)] \\ &\geq (r, \bar{u}) + (\varphi_x(x^*, u^*), y). \end{aligned} \quad (2.68a)$$

Then by the definition of the subgradient (2.62), we find that

$$(\varphi_x(x^*, u^*); r) \subseteq \partial_{x,u} \varphi(x^*, u^*).$$

The concavity case is very similar as we see that if we note that $(-\varphi)$ is convex. The lemma is proven. $\square$

Now we are ready to formulate the central result of this section.

Theorem 2.5 (Sufficient condition of optimality) *Let, under the assumptions (A1)–(A3) and (A5), the pair $(x^*(\cdot), u^*(\cdot))$ be an admissible pair and $\psi(t)$ be the corresponding adjoint variable satisfying (2.29). Assume that $h_0(x)$ and $g_l(x)$ ($l = \overline{1, L}$) are convex and $H(\psi(t), x, u, t)$ (2.31) is concave in (x, u) for any fixed $t \in [0, T]$ and any $\psi(t) \in \mathbb{R}^n$. Then this pair $(x^*(\cdot), u^*(\cdot))$ is optimal in the sense that the cost functional obeys $J(u(\cdot)) = h_0(x(T))$ (2.5) if*

$$H(\psi(t), x^*(t), u^*(t), t) = \max_{u \in U} H(\psi(t), x^*(t), u, t) \quad (2.69)$$

at almost all $t \in [0, T]$.

Proof By (2.69) and in view of the criterion of optimality (2.63), it follows that

$$0 \in \partial_u H(\psi(t), x^*(t), u^*(t), t). \quad (2.70)$$

Then, by the concavity of $H(\psi(t), x, u, t)$ in (x, u) , for any admissible pair (x, u) , and applying the integration operation, in view of (2.70), we get

$$\begin{aligned} & \int_{t=0}^T H(\psi(t), x(t), u(t), t) dt - \int_{t=0}^T H(\psi(t), x^*(t), u^*(t), t) dt \\ & \leq \int_{t=0}^T \left[\left(\frac{\partial}{\partial x} H(\psi(t), x^*(t), u^*(t), t), x(t) - x^*(t) \right) + (0, u(t) - u^*(t)) \right] dt \\ & = \int_{t=0}^T \left(\frac{\partial}{\partial x} H(\psi(t), x^*(t), u^*(t), t), x(t) - x^*(t) \right) dt. \end{aligned} \quad (2.71)$$

By the same trick as previously, let us introduce the “sensitivity” process $\delta(t) := x(t) - x^*(t)$, which evidently satisfies

$$\begin{aligned} \dot{\delta}(t) &= \eta(t) \quad \text{a.e. } t \in [0, T], \\ \delta(0) &= 0, \end{aligned} \quad (2.72)$$

where

$$\eta(t) := f(x(t), u(t), t) - f(x^*(t), u^*(t), t). \quad (2.73)$$

Then, in view of (2.29) and (2.71), it follows that

$$\begin{aligned} \frac{\partial}{\partial x} h_0(x^*(T))^T \delta(T) &= -[\psi(T)^T \delta(T) - \psi(0)^T \delta(0)] \\ &= -\int_{t=0}^T d[\psi(t)^T \delta(t)] \\ &= \int_{t=0}^T \frac{\partial}{\partial x} H(\psi(t), x^*(t), u^*(t), t)^T \delta(t) dt \end{aligned}$$

$$\begin{aligned}
& - \int_{t=0}^T \psi(t)^T (f(x(t), u(t), t) - f(x^*(t), u^*(t), t)) dt \\
& \geq \int_{t=0}^T [H(\psi(t), x(t), u, t) - H(\psi(t), x^*(t), u^*, t)] dt \\
& - \int_{t=0}^T \psi(t)^T (f(x(t), u(t), t) - f(x^*(t), u^*(t), t)) dt = 0.
\end{aligned} \tag{2.74}$$

The convexity of $h_0(x)$ and $g_l(x)$ ($l = 1, \dots, L$) and the complementary slackness condition yield

$$\begin{aligned}
& \left[\frac{\partial}{\partial x} h_0(x^*(T)) + \sum_{l=1}^L v_l \frac{\partial}{\partial x} g_l(x^*(T)) \right]^T \delta(T) \\
& \leq h_0(x(T)) - h_0(x^*(T)) + \sum_{l=1}^L v_l g_l(x^*(T)) \\
& = h_0(x(T)) - h_0(x^*(T)).
\end{aligned} \tag{2.75}$$

Combining (2.74) with (2.75), we derive

$$J(u(\cdot)) - J(u^*(\cdot)) = h_0(x(T)) - h_0(x^*(T)) \geq 0$$

and, since $u(\cdot)$ is arbitrary, the desired result follows. $\square$

Remark 2.2 Notice that checking the concavity property of $H(\psi(t), x, u, t)$ (2.31) in (x, u) for any fixed $t \in [0, T]$ and any $\psi(t) \in \mathbb{R}^n$ is not a simple task since it depends on the sign of the $\psi_i(t)$ components. So, the theorem given earlier may be applied directly practically only for a very narrow class of particular problems where the concavity property may be analytically checked.

2.3 Appendix

2.3.1 Linear ODE and Liouville's Theorem

Lemma 2.4 *The solution $x(t)$ of the linear ODE*

$$\begin{aligned}
\dot{x}(t) &= A(t)x(t), \quad t \geq t_0, \\
x(t_0) &= x_0 \in \mathbb{R}^{n \times n},
\end{aligned} \tag{2.76}$$

where $A(t)$ is an (almost everywhere) measurable matrix function, may be presented as

$$x(t) = \Phi(t, t_0)x_0, \tag{2.77}$$

where the matrix $\Phi(t, t_0)$ is the so-called fundamental matrix of the system (2.76) and satisfies the matrix ODE

$$\begin{aligned} \frac{d}{dt} \Phi(t, t_0) &= A(t) \Phi(t, t_0), \\ \Phi(t_0, t_0) &= I \end{aligned} \quad (2.78)$$

and verifies the group property

$$\Phi(t, t_0) = \Phi(t, s) \Phi(s, t_0) \quad \forall s \in (t_0, t). \quad (2.79)$$

Proof Assuming (2.77), direct differentiation of (2.77) implies

$$\dot{x}(t) = \frac{d}{dt} \Phi(t, t_0) x_0 = A(t) \Phi(t, t_0) x_0 = A(t) x(t).$$

So (2.77) verifies (2.76). The property (2.79) follows from the fact that

$$x(t) = \Phi(t, s) x_s = \Phi(t, s) \Phi(s, t_0) x_{t_0} = \Phi(t, t_0) x_{t_0}. \quad \square$$

Theorem 2.6 (Liouville) *If $\Phi(t, t_0)$ is the solution to (2.78), then*

$$\det \Phi(t, t_0) = \exp \left\{ \int_{s=t_0}^t \operatorname{tr} A(s) ds \right\}. \quad (2.80)$$

Proof The usual expansion for the determinant $\det \Phi(t, t_0)$ and the rule for differentiating the product of scalar functions show that

$$\frac{d}{dt} \det \Phi(t, t_0) = \sum_{j=1}^n \det \tilde{\Phi}_j(t, t_0),$$

where $\tilde{\Phi}_j(t, t_0)$ is the matrix obtained by replacing the j th row

$$\Phi_{j,1}(t, t_0), \quad \dots, \quad \Phi_{j,n}(t, t_0)$$

of $\Phi(t, t_0)$ by its derivatives

$$\dot{\Phi}_{j,1}(t, t_0), \quad \dots, \quad \dot{\Phi}_{j,n}(t, t_0).$$

But since

$$\dot{\Phi}_{j,k}(t, t_0) = \sum_{i=1}^n a_{j,i}(t) \Phi_{i,k}(t, t_0),$$

$$A(t) = \|a_{j,i}(t)\|_{j,i=1,\dots,n}$$

it follows that

$$\det \tilde{\Phi}_j(t, t_0) = a_{j,j}(t) \det \Phi(t, t_0),$$

which gives

$$\begin{aligned}\frac{d}{dt} \det \Phi(t, t_0) &= \sum_{j=1}^n \frac{d}{dt} \det \tilde{\Phi}_j(t, t_0) = \sum_{j=1}^n a_{j,j}(t) \det \Phi(t, t_0) \\ &= \operatorname{tr}\{A(t)\} \det \Phi(t, t_0)\end{aligned}$$

and, as a result, we obtain (2.80). $\square$

Corollary 2.5 *If for the system (2.76)*

$$\int_{s=t_0}^T \operatorname{tr} A(s) ds > -\infty \quad (2.81)$$

then for any $t \in [t_0, T]$

$$\det \Phi(t, t_0) > 0. \quad (2.82)$$

Proof It is the direct consequence of (2.80). $\square$

Lemma 2.5 *If*

$$\int_{s=t_0}^T \operatorname{tr} A(s) ds > -\infty$$

then the solution $x(t)$ of the linear nonautonomous ODE

$$\begin{aligned}\dot{x}(t) &= A(t)x(t) + f(t), \quad t \geq t_0, \\ x(t_0) &= x_0 \in \mathbb{R}^{n \times n},\end{aligned} \quad (2.83)$$

where $A(t)$ and $f(t)$ are assumed to be (almost everywhere) measurable matrix and vector functions, may be represented as (this is the Cauchy formula)

$$x(t) = \Phi(t, t_0) \left[x_0 + \int_{s=t_0}^t \Phi^{-1}(s, t_0) f(s) ds \right], \quad (2.84)$$

where $\Phi^{-1}(t, t_0)$ exists for all $t \in [t_0, T]$ and satisfies

$$\begin{aligned}\frac{d}{dt} \Phi^{-1}(t, t_0) &= -\Phi^{-1}(t, t_0) A(t), \\ \Phi^{-1}(t_0, t_0) &= I.\end{aligned} \quad (2.85)$$

Proof By the previous corollary, $\Phi^{-1}(t, t_0)$ exists within the interval $[t_0, T]$. Direct derivation of (2.84) implies

$$\dot{x}(t) = \dot{\Phi}(t, t_0) \left[x_0 + \int_{s=t_0}^t \Phi^{-1}(s, t_0) f(s) ds \right] + \Phi(t, t_0) \Phi^{-1}(t, t_0) f(t)$$

$$\begin{aligned}
&= A(t)\Phi(t, t_0) \left[x_0 + \int_{s=t_0}^t \Phi^{-1}(s, t_0) f(s) ds \right] + f(t) \\
&= A(t)x(t) + f(t),
\end{aligned}$$

which coincides with (2.83). Notice that the integral in (2.84) is well defined in view of the measurability property of the participating functions to be integrated. By the identities

$$\begin{aligned}
\Phi(t, t_0)\Phi^{-1}(t, t_0) &= I, \\
\frac{d}{dt} [\Phi(t, t_0)\Phi^{-1}(t, t_0)] &= \dot{\Phi}(t, t_0)\Phi^{-1}(t, t_0) + \Phi(t, t_0)\frac{d}{dt}\Phi^{-1}(t, t_0) = 0,
\end{aligned}$$

it follows that

$$\begin{aligned}
\frac{d}{dt}\Phi^{-1}(t, t_0) &= -\Phi^{-1}(t, t_0)[\dot{\Phi}(t, t_0)]\Phi^{-1}(t, t_0) \\
&= -\Phi^{-1}(t, t_0)[A(t)\Phi(t, t_0)]\Phi^{-1}(t, t_0) = -\Phi^{-1}(t, t_0)A(t).
\end{aligned}$$

The lemma is proven. $\square$

Remark 2.3 The solution (2.84) can be rewritten as

$$x(t) = \Phi(t, t_0)x_0 + \int_{s=t_0}^t \Phi(t, s)f(s) ds \quad (2.86)$$

since by (2.79)

$$\Phi(t, s) = \Phi(t, t_0)\Phi^{-1}(s, t_0).$$

2.3.2 Bihari Lemma

Lemma 2.6 (Bihari) *Let*

(1) *$v(t)$ and $\xi(t)$ be nonnegative continuous functions on $[t_0, \infty)$, that is,*

$$v(t) \geq 0, \quad \xi(t) \geq 0 \quad \forall t \in [t_0, \infty), \quad v(t), \xi(t) \in C[t_0, \infty). \quad (2.87)$$

(2) *For any $t \in [t_0, \infty)$ the following inequality holds:*

$$v(t) \leq c + \int_{\tau=t_0}^t \xi(\tau)\Phi(v(\tau)) d\tau, \quad (2.88)$$

where c is a positive constant ($c > 0$) and $\Phi(v)$ is a positive nondecreasing continuous function, that is,

$$0 < \Phi(v) \in C[t_0, \infty), \quad \forall v \in (0, \bar{v}), \quad \bar{v} \leq \infty. \quad (2.89)$$

Denote

$$\Psi(v) := \int_{s=c}^v \frac{ds}{\Phi(s)} \quad (0 < v < \bar{v}). \quad (2.90)$$

If, in addition,

$$\int_{\tau=t_0}^t \xi(\tau) d\tau < \Psi(\bar{v} - 0), \quad t \in [t_0, \infty) \quad (2.91)$$

then for any $t \in [t_0, \infty)$

$$v(t) \leq \Psi^{-1} \left(\int_{\tau=t_0}^t \xi(\tau) d\tau \right), \quad (2.92)$$

where $\Psi^{-1}(y)$ is the function inverse to $\Psi(v)$, that is,

$$y = \Psi(v), \quad v = \Psi^{-1}(y). \quad (2.93)$$

In particular, if $\bar{v} = \infty$ and $\Psi(\infty) = \infty$, then the inequality (2.92) is fulfilled without any constraints.

Proof Since $\Phi(v)$ is a positive nondecreasing continuous function the inequality (2.88) implies that

$$\Phi(v(t)) \leq \Phi \left(c + \int_{\tau=t_0}^t \xi(\tau) \Phi(v(\tau)) d\tau \right)$$

and

$$\frac{\xi(t) \Phi(v(t))}{\Phi(c + \int_{\tau=t_0}^t \xi(\tau) \Phi(v(\tau)) d\tau)} \leq \xi(t).$$

Integrating the last inequality, we obtain

$$\int_{s=t_0}^t \frac{\xi(s) \Phi(v(s))}{\Phi(c + \int_{\tau=t_0}^s \xi(\tau) \Phi(v(\tau)) d\tau)} ds \leq \int_{s=t_0}^t \xi(s) ds. \quad (2.94)$$

Denote

$$w(t) := c + \int_{\tau=t_0}^t \xi(\tau) \Phi(v(\tau)) d\tau.$$

Then evidently

$$\dot{w}(t) = \xi(t) \Phi(v(t)).$$

Hence, in view of (2.90), the inequality (2.94) may be represented as

$$\int_{s=t_0}^t \frac{\dot{w}(s)}{\Phi(w(s))} ds = \int_{w=w(t_0)}^{w(t)} \frac{dw}{\Phi(w)} = \Psi(w(t)) - \Psi(w(t_0)) \leq \int_{s=t_0}^t \xi(s) ds.$$

Taking into account that $w(t_0) = c$ and $\Psi(w(t_0)) = 0$, from the last inequality it follows that

$$\Psi(w(t)) \leq \int_{s=t_0}^t \xi(s) ds. \quad (2.95)$$

Since

$$\Psi'(v) = \frac{1}{\Phi(v)} \quad (0 < v < \bar{v})$$

the function $\Psi(v)$ has the uniquely defined, continuous monotonically increasing inverse function $\Psi^{-1}(y)$ given within the open interval $(\Psi(+0), \Psi(\bar{v}-0))$. Hence, (2.95) directly implies

$$w(t) = c + \int_{\tau=t_0}^t \xi(\tau) \Phi(v(\tau)) d\tau \leq \Psi^{-1} \left(\int_{s=t_0}^t \xi(s) ds \right),$$

which, in view of (2.88), leads to (2.92). Indeed,

$$v(t) \leq c + \int_{\tau=t_0}^t \xi(\tau) \Phi(v(\tau)) d\tau \leq \Psi^{-1} \left(\int_{s=t_0}^t \xi(s) ds \right).$$

The case of $\bar{v} = \infty$ and $\Psi(\infty) = \infty$ is evident. The lemma is proven. $\square$

Corollary 2.6 Taking in (2.92)

$$\Phi(v) = v^m \quad (m > 0, m \neq 1),$$

it follows that

$$\boxed{v(t) \leq \left[c^{1-m} + (1-m) \int_{\tau=t_0}^t \xi(\tau) d\tau \right]^{\frac{1}{m-1}} \quad \text{for } 0 < m < 1} \quad (2.96)$$

and

$$\boxed{v(t) \leq c \left[1 - (1-m)c^{m-1} \int_{\tau=t_0}^t \xi(\tau) d\tau \right]^{-\frac{1}{m-1}} \quad \text{for } m > 1 \text{ and } \int_{\tau=t_0}^t \xi(\tau) d\tau < \frac{1}{(m-1)c^{m-1}}.}$$

2.3.3 Gronwall Lemma

Corollary 2.7 (Gronwall) *If $v(t)$ and $\xi(t)$ are nonnegative continuous functions in $[t_0, \infty)$ verifying*

$$v(t) \leq c + \int_{\tau=t_0}^t \xi(\tau) v(\tau) d\tau \quad (2.97)$$

then for any $t \in [t_0, \infty)$ the following inequality holds:

$$v(t) \leq c \exp\left(\int_{s=t_0}^t \xi(s) ds\right). \quad (2.98)$$

This result remains true if $c = 0$.

Proof Taking in (2.88) and (2.90)

$$\Phi(v) = v,$$

we obtain (2.97) and, hence, for the case $c > 0$

$$\Psi(v) := \int_{s=c}^v \frac{ds}{s} = \ln\left(\frac{v}{c}\right)$$

and

$$\Psi^{-1}(y) = c \cdot \exp(y),$$

which implies (2.98). The case $c = 0$ follows from (2.98) on applying $c \rightarrow 0$. $\square$

2.3.4 The Lagrange Principle in Finite-Dimensional Spaces

Let us recall several simple and commonly used definitions.

Definition 2.7 A set C lying within a linear set X is called *convex* if, together with any two points $x, y \in C$, it also contains the closed interval

$$[x, y] := \{z : z = \alpha x + (1 - \alpha)y, \alpha \in [0, 1]\}. \quad (2.99)$$

A function $f : X \rightarrow \mathbb{R}$ is said to be *convex* if for any $x, y \in X$ and any $\alpha \in [0, 1]$

$$f(\alpha x + (1 - \alpha)y) \leq \alpha f(x) + (1 - \alpha)f(y) \quad (2.100)$$

or, in other words, if the *supergraph* of f defined as

$$\text{epi } f = \{(a, x) \in \mathbb{R} \times X : a \geq f(x)\} \quad (2.101)$$

is a convex set in $\mathbb{R} \times X$.

Consider the optimal control problem in the Mayer form (2.5), that is,

$$\begin{aligned} J(u(\cdot)) = h_0(x(T)) &\rightarrow \min_{u(\cdot) \in \mathcal{U}_{\text{admis}}[0, T]}, \\ x(T) \in \mathcal{M} &= \{x \in \mathbb{R}^n : g_l(x) \leq 0 \ (l = \overline{1, L})\}, \end{aligned} \quad (2.102)$$

where $h_0(x)$ and $g_l(x)$ ($l = 1, \dots, L$) are *convex functions*. For the corresponding optimal pair $(x^*(\cdot), u^*(t))$ it follows that

$$\begin{aligned} J(u^*(\cdot)) = h_0(x^*(T)) &\leq J(u(\cdot)) = h_0(x(T)), \\ g_l(x^*(T)) &\leq 0 \quad (l = 1, \dots, L) \end{aligned} \quad (2.103)$$

for any $u(t)$ and corresponding $x(T)$ satisfying

$$g_l(x(T)) \leq 0 \quad (l = 1, \dots, L). \quad (2.104)$$

Theorem 2.7 (The Lagrange Principle, Kuhn and Tucker 1951) *Let X be a linear (not necessarily finite-dimensional) space,*

$$h_0 : X \rightarrow \mathbb{R}$$

and let

$$g_l : X \rightarrow \mathbb{R} \quad (l = 1, \dots, L)$$

be convex functions in X and X_0 be a convex subset of X , that is, $X_0 \in X$.

A. If $(x^*(\cdot), u^*(t))$ is an optimal pair then there exist nonnegative constants $\mu^* \geq 0$ and $v_l^* \geq 0$ ($l = 1, \dots, L$) such that the following two conditions hold.

(1) “Minimality condition for the Lagrange function”:

$$\begin{aligned} L(x^*(T), \mu^*, v^*) &\leq L(x(T), \mu^*, v^*), \\ L(x(T), \mu, v) &:= \mu^* h_0(x(T)) + \sum_{l=1}^L v_l^* g_l(x(T)), \end{aligned} \quad (2.105)$$

where $x(T)$ corresponds to any admissible control $u(\cdot)$.

(2) “Complementary slackness”:

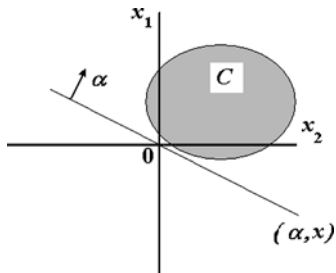
$$v_l^* g_l(x^*(T)) = 0 \quad (l = 1, \dots, L). \quad (2.106)$$

B. If $\mu^* > 0$ (the regular case), then the conditions (1)–(2) turn out to be sufficient to guarantee that $(x^*(\cdot), u^*(t))$ is a solution of the problem (2.102).

C. To guarantee that there exists $\mu^* > 0$, it is sufficient that the so-called Slater condition holds, that is, there exists an admissible pair

$$(\bar{x}(\cdot), \bar{u}(t)) \quad (\bar{x}(T) \in \mathcal{M})$$

Fig. 2.1 The illustration of the *Separation Principle*



such that

$$g_l(\bar{x}(T)) < 0 \quad (l = 1, \dots, L). \quad (2.107)$$

The considerations in the following follow from Alexeev et al. (1979).

The Separation Principle

First, we will formulate and prove the theorem called *the Separation Principle for a finite-dimensional space*, which plays a key role in the proof of the Lagrange Principle.

Theorem 2.8 (The Separation Principle) *Let $C \subseteq \mathbb{R}^n$ be a convex subspace of $\mathbb{R}^n$ which does not contain the point 0, that is, $0 \notin C$. Then there exists a vector $\alpha = (\alpha_1, \dots, \alpha_n)^T \in \mathbb{R}^n$ such that for any $x = (x_1, \dots, x_n)^T \in C$ the following inequality holds:*

$$\sum_{i=1}^n \alpha_i x_i \geq 0. \quad (2.108)$$

In other words, the plane

$$\sum_{i=1}^n \alpha_i x_i = 0$$

separates the space $\mathbb{R}^n$ in two subspaces, one of which contains the set C completely (see Fig. 2.1).

Proof Let $\text{lin } C$ be a minimal linear subspace of $\mathbb{R}^n$ containing C . Only two cases are possible

$$\text{lin } C \neq \mathbb{R}^n \quad \text{or} \quad \text{lin } C = \mathbb{R}^n.$$

1. If $\text{lin } C \neq \mathbb{R}^n$, then $\text{lin } C$ is a proper subspace in $\mathbb{R}^n$ and, therefore, there exists a hyperplane

$$\sum_{i=1}^n \alpha_i x_i = 0$$

containing C as well as the point 0. This plane may be selected as the one we are interested in.

2. If $\text{lin } C = \mathbb{R}^n$, then from the vectors belonging to C we may select n linearly independent ones forming a basis in $\mathbb{R}^n$. Denote them by

$$e^1, \dots, e^n \quad (e^i \in C, i = 1, \dots, n).$$

Consider then the two convex sets (more exactly, cones): a nonnegative “orthant” K_1 and a “convex cone” K_2 , defined as

$$\begin{aligned} K_1 &:= \left\{ x \in \mathbb{R}^n : x = \sum_{i=1}^n \beta_i e^i, \beta_i \geq 0 \right\}, \\ K_2 &:= \left\{ x \in \mathbb{R}^n : x = \sum_{i=1}^s \alpha_i \bar{e}^i, \alpha_i \geq 0, \bar{e}^i \in C, \right. \\ &\quad \left. i = 1, \dots, s \text{ (} s \in \mathbb{N} \text{ is any natural number)} \right\}. \end{aligned} \quad (2.109)$$

These two cones are not crossed, that is, they do not contain a common point. Indeed, suppose that there exists a vector

$$\bar{x} = - \sum_{i=1}^n \bar{\beta}_i e^i, \quad \bar{\beta}_i > 0$$

which also belongs to K_2 . Then necessarily one finds $s \in \mathbb{N}$, $\bar{\alpha}_i \geq 0$ and $\bar{e}^i$ such that

$$\bar{x} = \sum_{i=1}^s \bar{\alpha}_i \bar{e}^i.$$

But this is possible only if $0 \in C$ since, in this case, the point 0 might be represented as a convex combination of some points from C , that is,

$$\begin{aligned} 0 &= \frac{\sum_{i=1}^s \bar{\alpha}_i \bar{e}^i - \bar{x}}{\sum_{i=1}^s \bar{\alpha}_i + \sum_{i=1}^n \bar{\beta}_i} = \frac{\sum_{i=1}^s \bar{\alpha}_i \bar{e}^i + \sum_{i=1}^n \bar{\beta}_i e^i}{\sum_{i=1}^s \bar{\alpha}_i + \sum_{i=1}^n \bar{\beta}_i} \\ &= \sum_{i=1}^s \frac{(\bar{\alpha}_i + \bar{\beta}_i)}{\sum_{j=1}^s (\bar{\alpha}_j + \bar{\beta}_j)} e^i. \end{aligned} \quad (2.110)$$

But this contradicts the assumption that $0 \notin C$. So,

$$K_1 \cap K_2 = \emptyset. \quad (2.111)$$

3. Since $\mathcal{K}_1$ is an open set, any point $x \in \mathcal{K}_1$ cannot belong to $\bar{\mathcal{K}}_2$ at the same time. Here $\bar{\mathcal{K}}_2$ is the “closure” of $\mathcal{K}_2$. Note that $\bar{\mathcal{K}}_2$ is a closed and convex set. Let us consider any point $x^0 \in \mathcal{K}_1$, for example,

$$x^0 = - \sum_{i=1}^n \tilde{e}^i,$$

and try to find a point $y^0 \in \bar{\mathcal{K}}_2$ closer to x^0 . Such a point necessarily exists; namely, it is the point which minimizes the continuous function f

$$(y) := \|x - y\|$$

within all y belonging to the compact set

$$\bar{\mathcal{K}}_2 \cap \{x \in \mathcal{K}_1 : \|x - x^0\| \leq \varepsilon \text{ small enough}\}.$$

4. Then let us construct the hyperplane H orthogonal to $(x^0 - y^0)$ and show that this is the plane that we are interested in, that is, show that $0 \in H$ and C belongs to a half-closed subspace separated by this surface, namely,

$$(\text{int } H \cap \bar{\mathcal{K}}_2) = \emptyset.$$

Also, since $C \subseteq \bar{\mathcal{K}}_2$, we have

$$C \subseteq \overline{(\mathbb{R}^n \setminus \text{int } H)}.$$

By contradiction, let us suppose that there exists a point

$$\tilde{y} \in (\text{int } H \cap \bar{\mathcal{K}}_2).$$

Then the angle $\angle x^0 y^0 \tilde{y}$ is less than $\pi/2$, and, besides, since $\bar{\mathcal{K}}_2$ is convex, it follows that

$$[y^0, \tilde{y}] \in \bar{\mathcal{K}}_2.$$

Let us take the point

$$\tilde{y}' \in (y^0, \tilde{y})$$

such that

$$(x^0, \tilde{y}') \perp (y^0, \tilde{y})$$

and show that $\tilde{y}'$ is not a point of $\bar{\mathcal{K}}_2$ close to x^0 . Indeed, the points y^0 , $\tilde{y}$, and $\tilde{y}'$ belong to the same line and $\tilde{y}' \in \text{int } H$. But if

$$\tilde{y}' \in [y^0, \tilde{y}]$$

and $\tilde{y}' \in \bar{\mathcal{K}}_2$, then it is necessary that we have

$$\|x^0 - \tilde{y}'\| < \|x^0 - y^0\|$$

(the shortest distance is one smaller than any other one). At the same time, $\tilde{y}' \in (y^0, \tilde{y})$, so

$$\|x^0 - \tilde{y}\| < \|x^0 - y^0\|.$$

Also we have $0 \in H$ since, if this were not so, the line $[0, \infty)$, crossing y^0 and belonging to $\tilde{K}_2$, should necessarily have points in common with $\text{int } H$. $\square$

Proof of the Lagrange Principle

Now we are ready to give the proof of the main claim.

Proof of the Lagrange Principle Let $(x^*(\cdot), u^*(t))$ be an optimal pair. For a new normalized cost function defined as

$$\tilde{J}(u(\cdot)) := J(u(\cdot)) - J(u^*(\cdot)) = h_0(x(T)) - h_0(x^*(T)) \quad (2.112)$$

it follows that

$$\min_{u(\cdot) \in \mathcal{U}_{\text{admis}}} \tilde{J}(u(\cdot)) = 0. \quad (2.113)$$

Define

$$C := \left\{ \eta \in \mathbb{R}^{L+1} \mid \exists x \in X_0: h_0(x) - h_0(x^*(T)) < \eta_0, \right. \\ \left. g_l(x) \leq \eta_l \ (l = 1, \dots, L) \right\}. \quad (2.114)$$

(A) *The set C is nonempty and convex.* Indeed, the vector η with positive components belongs to C since in (2.114) we may take $x = x^*(T)$. Hence C is nonempty. Let us prove its convexity. Consider two vectors η and η' both belonging to C , $\alpha \in [0, 1]$ and $x, x' \in X_0$ such that for all $l = 1, \dots, L$

$$\begin{aligned} h_0(x) - h_0(x^*(T)) &< \eta_0, & g_l(x) &\leq \eta_l, \\ h_0(x') - h_0(x^*(T)) &< \eta'_0, & g_l(x') &\leq \eta'_l. \end{aligned} \quad (2.115)$$

Denote

$$x^\alpha := \alpha x + (1 - \alpha)x'.$$

In view of the convexity of X_0 it follows that $x^\alpha \in X_0$. The convexity of the functions $h_0(x)$ and $g_l(x)$ ($l = 1, \dots, L$) implies that

$$\alpha \eta + (1 - \alpha) \eta' \in C. \quad (2.116)$$

Indeed,

$$\begin{aligned} h_0(x^\alpha) - h_0(x^*(T)) \\ \leq \alpha [h_0(x) - h_0(x^*(T))] + (1 - \alpha) [h_0(x') - h_0(x^*(T))] \end{aligned}$$

$$\begin{aligned} &\leq \alpha\eta_0 + (1 - \alpha)\eta'_0, \\ g_l(x^\alpha) &\leq \alpha g_l(x) + (1 - \alpha)g_l(x') \leq \alpha\eta_l + (1 - \alpha)\eta'_l \quad (l = 1, \dots, L). \end{aligned}$$

So C is nonempty and convex.

(B) *The point 0 does not belong to C .* Indeed, if it did, in view of the definition (2.114), there would exist a point $x \in X_0$ satisfying

$$\begin{aligned} h_0(x^\alpha) - h_0(x^*(T)) &< 0, \\ g_l(x^\alpha) &\leq 0 \quad (l = 1, \dots, L), \end{aligned} \tag{2.117}$$

which is in contradiction to the fact that $x^*(T)$ is a solution of the problem. So $0 \notin C$. In view of this fact and taking into account the convexity property of C , we may apply the Separation Principle for a finite-dimensional space: there exist constants

$$\mu^*, v_1^*, \dots, v_L^*$$

such that for all $\eta \in C$

$$\mu^*\eta_0 + \sum_{l=1}^L v_l^*\eta_l \geq 0. \tag{2.118}$$

(C) *The multipliers μ^* and v_l^* ($l = 1, \dots, L$) in (2.118) are nonnegative.* In (A) we have already mentioned that any vector $\eta \in \mathbb{R}^{L+1}$ with positive components belongs to C , and, in particular, the vector

$$\left(\underbrace{\varepsilon, \dots, \varepsilon}_{l_0}, 1, \varepsilon, \dots, \varepsilon \right)$$

($\varepsilon > 0$) does. The substitution of this vector into (2.118) leads to the following inequalities:

$$\begin{cases} v_{l_0}^* \geq -\mu^*\varepsilon - \varepsilon \sum_{l=l_0}^L v_l^* & \text{if } 1 \leq l_0 \leq L, \\ \mu^* \geq -\varepsilon \sum_{l=1}^L v_l^* & \text{if } l_0 = 0. \end{cases} \tag{2.119}$$

Letting ε go to zero in (2.119) implies the nonnegativity property for the multipliers μ^* and v_l^* ($l = 1, \dots, L$).

(D) *The multipliers v_l^* ($l = 1, \dots, L$) satisfy the complementary slackness condition (2.106).* Indeed, if

$$g_{l_0}(x^*(T)) = 0$$

then the identity

$$v_{l_0}^* g_{l_0}(x^*(T)) = 0$$

is trivial. Suppose that

$$g_{l_0}(x^*(T)) < 0.$$

Then for $\delta > 0$ the point

$$\left(\underbrace{\delta, 0, \dots, 0, g_{l_0}(x^*(T))}_{l_0}, 0, \dots, 0 \right) \quad (2.120)$$

belongs to the set C . To check this, it is sufficient to take $x = x^*(T)$ in (2.118). The substitution of this point into (2.118) implies

$$v_{l_0}^* g_{l_0}(x^*(T)) \geq -\mu^* \delta. \quad (2.121)$$

Letting δ go to zero we obtain

$$v_{l_0}^* g_{l_0}(x^*(T)) \geq 0$$

and since

$$g_{l_0}(x^*(T)) < 0$$

it follows that

$$v_{l_0}^* \leq 0.$$

But in (C) it has been proven that $v_{l_0}^* \geq 0$. Thus, $v_{l_0}^* = 0$, and, hence,

$$v_{l_0}^* g_{l_0}(x^*(T)) = 0.$$

(E) *Minimality condition for the Lagrange function.* Let $x(T) \in X_0$. Then, as follows from (2.114), the point

$$([h_0(x(T)) - h_0(x^*(T))] + \delta, g_1(x(T)), \dots, g_L(x(T)))$$

belongs to C for any $\delta > 0$. The substitution of this point into (2.118), in view of (D), yields

$$\begin{aligned} L(x(T), \mu, v) &:= \mu^* h_0(x(T)) + \sum_{l=1}^L v_l^* g_l(x(T)) \\ &\geq \mu^* h_0(x^*(T)) - \mu^* \delta \\ &= \mu^* h_0(x^*(T)) + \sum_{l=1}^L v_l^* g_l(x^*(T)) - \mu^* \delta \\ &= L(x^*(T), \mu^*, v^*) - \mu^* \delta. \end{aligned} \quad (2.122)$$

Taking $\delta \rightarrow 0$ we obtain (2.105).

(F) If $\mu^* > 0$ (the regular case), then the conditions (A1) and (A2) of Theorem 2.7 are sufficient for optimality. Indeed, in this case it is clear that we may take $\mu^* = 1$, and, hence,

$$\begin{aligned} h_0(x(T)) &\geq h_0(x(T)) + \sum_{l=1}^L v_l^* g_l(x(T)) \\ &= L(x(T), 1, v^*) \geq L(x^*(T), 1, v^*) \\ &= h_0(x^*(T)) + \sum_{l=1}^L v_l^* g_l(x^*(T)) = h_0(x^*(T)). \end{aligned}$$

This means that $x^*(T)$ corresponds to an optimal solution.

(G) Slater's condition of regularity. Suppose that Slater's condition is fulfilled, but $\mu = 0$. We directly obtain a contradiction. Indeed, since not all v_l^* are equal to zero simultaneously, it follows that

$$L(\bar{x}(T), 0, v^*) = \sum_{l=1}^L v_l^* g_l(\bar{x}(T)) < 0 = L(x^*(T), 0, v^*),$$

which is in contradiction to (E). □

Corollary 2.8 (The saddle-point property for the regular case) *In the regular case the so-called saddle-point property holds:*

$$L(x(T), 1, v^*) \geq L(x^*(T), 1, v^*) \geq L(x^*(T), 1, v) \quad (2.123)$$

or, in another form,

$$\min_{u(\cdot) \in \mathcal{U}_{\text{admis}}} L(x^*(T), 1, v^*) = L(x^*(T), 1, v^*) = \max_{v \geq 0} L(x^*(T), 1, v). \quad (2.124)$$

Proof The left-hand side inequality has been already proven in (F). As for the right-hand side inequality, it directly follows from

$$\begin{aligned} L(x^*(T), 1, v^*) &= h_0(x^*(T)) + \sum_{l=1}^L v_l^* g_l(x^*(T)) \\ &= h_0(x^*(T)) \geq h_0(x^*(T)) + \sum_{l=1}^L v_l g_l(x^*(T)) \\ &= L(x^*(T), 1, v). \end{aligned} \quad (2.125)$$

□

Remark 2.4 The construction of the Lagrange function in the form

$$L(x, \mu, v) = \mu h_0(x) + \sum_{l=1}^L v_l g_l(x) \quad (2.126)$$

with $\mu \geq 0$ is very essential since the use of this form only as $L(x, 1, v)$, when the regularity conditions are not valid, may provoke a serious error in the optimization process. The following *counterexample* demonstrates this effect. Consider the simple static optimization problem formulated as

$$\begin{cases} h_0(x) := x_1 \rightarrow \min_{x \in \mathbb{R}^2}, \\ g(x) := x_1^2 + x_2^2 \leq 0. \end{cases} \quad (2.127)$$

This problem evidently has the unique solution

$$x_1 = x_2 = 0.$$

But the direct use of the Lagrange Principle with $\mu = 1$ leads to the following contradiction:

$$\begin{cases} L(x, 1, v^*) = x_1 + v^*(x_1^2 + x_2^2) \rightarrow \min_{x \in \mathbb{R}^2} \\ \frac{\partial}{\partial x_1} L(x^*, 1, v^*) = 1 + 2v^*x_1^* = 0, \\ \frac{\partial}{\partial x_2} L(x^*, 1, v^*) = 2v^*x_2^* = 0, \\ v^* \neq 0, \quad x_2^* = 0, \quad x_1^* = -\frac{1}{2v^*} \neq 0. \end{cases} \quad (2.128)$$

Notice that for this example the Slater condition is not valid.



<http://www.springer.com/978-0-8176-8151-7>

The Robust Maximum Principle

Theory and Applications

Boltyanski, V.G.; Poznyak, A.

2012, XXII, 432 p. 36 illus., Hardcover

ISBN: 978-0-8176-8151-7

A product of Birkhäuser Basel